

An Introduction to Language Description

An Introduction to Language Description

By

John Rhoades

CAMBRIDGE
SCHOLARS

P U B L I S H I N G

An Introduction to Language Description
By John Rhoades

This book first published 2014

Cambridge Scholars Publishing

12 Back Chapman Street, Newcastle upon Tyne, NE6 2XX, UK

British Library Cataloguing in Publication Data
A catalogue record for this book is available from the British Library

Copyright © 2014 by John Rhoades

All rights for this book reserved. No part of this book may be reproduced, stored in a retrieval system, or transmitted, in any form or by any means, electronic, mechanical, photocopying, recording or otherwise, without the prior permission of the copyright owner.

ISBN (10): 1-4438-5357-7, ISBN (13): 978-1-4438-5357-6

TABLE OF CONTENTS

Figures	vii
Tables	viii
Preface	ix
Chapter One.....	1
Introduction	
Chapter Two	14
Phonetics: The Description of Speech	
Chapter Three	46
Phonology and Phonemic Analysis: Systems of Speech Sounds	
Chapter Four	75
Semantics: The Structure of Meaning	
Chapter Five	100
Grammar: The Organization of Words and Sentences	
Appendix A	124
Examples of Articulatory Configurations from Author's Dialect	
Appendix B	127
Example of a Phonetic Corpus	
Appendix C	131
The International Phonetic Alphabet (2005)	
Appendix D	132
English Inflectional and Derivational Affixes	

Bibliography	135
Exercises for Chapters Two–Five.....	137
Suggested Solutions to Exercises	156
Index	175

FIGURES

Figure 2-1 Spectrogram of Human Speech.....	16
Figure 2-2 Cross-Section of Human Vocal Tract.....	21
Figure 2-3 Simplified Phonetic Chart for Consonants Showing Place and Manner of Articulation.....	24
Figure 2-4 Simplified Phonetic Chart for Vowels Showing Tongue Height and Tongue Height Location.....	25
Figure 2-5 Expanded Consonant Chart.....	29
Figure 2-6 Expanded Vowel Charts.....	31
Figure 2-7 Forward Glides.....	33
Figure 2-8 Mid and Central Glides	33
Figure 2-9 Consonsant Chart with Notation Graphs.....	38
Figure 2-10 Vowel Charts with Notation Gaphs.....	39
Figure 3-1 Formula for Monosyllabic English Word.....	73
Figure 4-1 Primate Taxonomy	97

TABLES

Table 2–1 Descriptive Terms for Places of Consonant Articulation	22
Table 2–2 Common English Words with Selected Sounds Described in Articulatory Terms.....	34

PREFACE

The study of language has no peer in the human sciences. As our species' distinctive attribute, language is the gateway to the mind, its symbolic facilities, and, indeed, what it is to be human. This prospect of gaining access to the essence of human behavior in all of its manifold complexity, however exciting, was actually not what drew the author to the study of language. In fact it was his inability to master language as a subject. English grammar presented in secondary school was a series of puzzling parsing diagrams; Latin (offered as an academic treat) was a series of mostly erroneous struggles to translate seemingly abstruse passages; and German in college, with its attendant presentation of German speech sounds (as well as grammar) by an infinitely patient but perpetually disappointed professor, simply served to cement the realization that "language" was best avoided as too forbidding or too difficult for an otherwise reasonably competent student. This is, I suspect, not an uncommon orientation for many students; it was certainly not an auspicious beginning for someone who would wind up studying language as a profession.

However, the chance registration in the anthropological linguistics courses offered by Professors Harry Hoijer and, later, Mary Woodward and Jeannette Witucki provided the liberating appreciation that language was patterned! Not mysterious, seemingly random, collections of aphorisms easily confused and forgotten, but instead patterned activity organized by comprehensible sets of rules. And even more intriguing was the realization that these rules could be discovered by careful observation and the application of analytic procedures. For a young student, attracted by the promise of empirical investigation of human behavior, this was a powerful revelation. But above all, a part of the human experience that had seemed so foreign and out of reach, was made accessible. Language could be comprehended, and this is what the text seeks to impart to the reader.

This text is an attempt to introduce the basic systems of language and how these can be investigated and described. The systems of organizing speech sounds (phonology), meaning (lexical semantics), and grammar (morphology and syntax) are described along with methods of figuring these out for other languages. The goal is to have the reader appreciate (hopefully with the same excitement) that language is patterned and that these patterns can be discovered. To this end, problems using actual

language data are presented at the end so the reader can apply the methods described in the chapter.

Further, it is important to understand that language structures are not simply abstract puzzles divorced from actual human speech communication. Throughout each chapter references are made to humans speaking and choosing among structural elements. This will hopefully ground the presentation in the realities of speech, a reality that each reader shares as a language user.

I would like to express my deep appreciation to my colleague Zhiming Zhao (formerly of the State University of New York, College at Geneseo). Our many hours of discussion, and argument, about language topics were greatly stimulating and contributed to the value of the present work. Indeed, several parts of Chapters 1 and 2 are the direct result of suggestions and phrasing of Dr. Zhao. I would also wish to thank St. John Fisher College for two sabbatical leaves during which most of the writing was accomplished. Finally I dedicate this book to Mariana, Kimberly and Tom who supported the work of a sometimes preoccupied husband and father.

Rochester, NY
2013

CHAPTER ONE

INTRODUCTION

Linguistics is the study of human communication. This chapter will introduce the subject of linguistics by describing the skills and knowledge necessary to engage in the most common and important form of human communication—speech. Our example will be a simple conversation, and the skills needed to engage in it include not only the physiological skills of speaking (and hearing) but also, and more importantly, several mental skills of organizing speech sounds into appropriate sequences. In addition, there are the knowledge skills of how to converse with another person by selecting particular topics and styles of speaking. The chapter concludes with a discussion of the place of linguistics within the human sciences, in particular anthropology—the significance of the study of language for an understanding of humans, language and human evolution, language in relation to primate communicative capabilities, and a description of four linguistic sub-fields.

The Complexity of Ordinary Conversation

Imagine the following scene: two ordinary adult members of a community meet on the street and engage in a short conversation. They exchange greetings, discuss an upcoming community event, say goodbye and part. There is certainly nothing unusual about this—it occurs continuously in all communities as a common form of interaction. Yet the set of behaviors which take place during this event represent the most complex phenomena known to behavioral scientists. This may seem to be an overstatement. After all, each of you have performed a similar activity innumerable times and I am sure you would believe that it requires little or no skill. In fact, I suspect that often you have done this while you were engaged as well in doing something else—thinking about a subsequent activity, keeping a pile of books in your grasp, chewing gum. However, the ability to converse and do one or more activities at the same time is not proof of the simplicity of conversation, but rather proof of the considerable and unequalled complexity of human cognitive and motor capabilities. To demonstrate this, let's consider all that our two adults must be able to do in order to adequately conduct their conversational interaction.

The Importance of Mental Skills

Before we start, however, we must examine a few basic assumptions. We will be unable to observe most of the abilities to be discussed. Their presence can only be inferred by other evidence than by just watching the conversation take place. These abilities are “knowledge” skills (such as the capability to use values, beliefs and meanings) and are considered to be part of the mind—the cognitive operations of the brain.

These are crucial abilities. Without them, the physiological skills of speaking and hearing would be insufficient for our two adults to converse. If we were to replace one of the adults with another, of equal physiology but of different cognitive abilities (i.e. of a different culture), then the conversation would not be adequately conducted to either adult’s satisfaction. They would simply not be able to greet each other, discuss something, or part in a manner that would be mutually intelligible. Another indication of the primacy of cognition is that humans are not restricted to speaking and hearing in order to communicate. Humans can also manifest their communicative signals in visual ways such as reading/writing and signing (a system of gestural signals used by the hearing impaired) and in tactile ways such as reading by Braille (raised dots felt by the fingertips). What is basic to human communication, then, is the cognitive capability to organize signals in a distinctively complex manner, not a particular mode of signaling. The scope of this book, however, will be limited to the mode of speaking (and hearing) since from an anthropological perspective speech is the basic mode of communication in human communities.

The Physiological Skills for Speaking: Phonetics

We will begin with the physiological abilities necessary for conversing. Speaking involves a very large number of physiological activities. Taken together, these activities are referred to as **articulation**. (And remember, we are only discussing an ordinary conversation, not some special tour-de-force of vocalization as might be produced by the lead singer in an especially energetic rock group, or the rapid-fire delivery of a professional auctioneer.) Articulation requires that the movements of hundreds of muscles, connected to many different organs and bones, be rapidly coordinated. In speech, the flow of air into and, especially, out of the lungs is continuously modified (in a manner, it should be noted, different from the rhythms of breathing). The flow of air from the lungs is further modified by rapidly changing configurations of the organs of the throat and mouth such as the lips and tongue. These are not only rapid and continuous, but multiple, which means that several configurations occur at

the same time in different portions of the vocal tract. One of the important occurring features of articulation is the “melody” of speaking—particular sequences of different pitch and loudness as well as patterns of pauses. The study of human articulation, including the methods of describing it in a useful manner, is called **phonetics**.

The Cognitive Skills for Speaking: Phonology

These complex articulatory events do not just happen as a result of the presence of the necessary physiological equipment, such as is the case for breathing or mastication and swallowing. They must be produced and organized by a cognitive system learned by the speaker. A speaker must learn to make certain kinds of sounds, rather than just any of the large number of sounds that are possible with the human vocal tract. Moreover, the speaker must learn to modify these sounds in particular ways when they are used in certain sound environments. Also, the speaker must learn how to use the sounds in combinations with other sounds. Not all combinations, however possible just on an articulatory basis, are permitted within the system of sound combinations learned by the speaker. The “melody” of speaking, in particular, will be learned such that a speaker produces certain variations of pitch and loudness rather than random variations. By the same reasoning, speakers will not employ a monotone of equally loud, pitched and evenly spaced units of sound (unless, of course, they have learned that this conveys a desired message to the hearer—among U. S. English speakers, for example, this might be boredom or restrained anger).

Linguists refer to the study of this cognitive system as **phonology**. We will take a certain terminological liberty for the sake of convenience and refer also to the cognitive system itself as “phonology.” The same fiction will be used for each of the cognitive systems described below. Instead of having to state, for example, that “linguists who have done a phonological study of the X group of speakers have come to the conclusion that these speakers must have a cognitive system which produces the following kinds of sounds: . . .,” we can simply state “the phonology of X language contains the following kinds of sounds: . . .”¹ Speakers possessing different phonologies will use their common human vocal tract organs to produce and use different kinds of speech sounds (although it should be noted that there will be many similarities among the world’s languages).

¹ Many linguists (the author included) use the term *phonemics*, or *phonemic analysis*, for the method of studying the cognitive system of sound production, and the term *phonology* for the system itself.

The Cognitive Skills for Speaking: Semantics

Sounds, however well-formed and acceptably combined, are really not the point of speaking—it is the use of these sounds to communicate. Sound groupings are signals which are normally produced to elicit a particular shared meaning in the mind of a hearer. Therefore, a further necessary cognitive skill is to be able to understand and use the meanings of the thousands of sound group signals required for ordinary adult conversation.

This should not be thought of as just a sort of mental dictionary with individual signals listed with their individual meanings. While these units of meaning may be used singly as complete conversational elements, their natural habitat is in combination with other such units and the result is rarely simply the sum of the meanings of the units but instead a function of the particular interplay among these constituent meanings. A speaker must be able to deal with the practically limitless, and often new, combinations of meanings that occur. And not just within the same sentence or phrase, but as constructed over a whole conversation and even many conversations. In addition, our speakers' abilities to use meaning units will depend upon understanding such relationships among meanings as synonymy, analogy, taxonomy, and the applicability of certain meanings with certain real world situations—i.e. relevance and truth. These skills of knowing and using meanings are referred to as **semantics** by linguists.

The Cognitive Skills for Speaking: Pragmatics

This last aspect mentioned above, the relationship between speaking and the external context of speaking, requires another ability. Our two adults do not possess just one way of speaking. They would, for example, undoubtedly have formal and informal manners of speaking (called **speech varieties**) and we could suppose that, as they are speaking in a public setting (a street), the formal variety might be more appropriate than the informal variety. The actual “rules” for variety selection would be a necessary cognitive skill for a speaker. There are typically several varieties available to adult speakers (from versions of the same language, to completely different languages, as in a bilingual community) and choosing a variety is related to such factors as the intentions of the speakers, their sex and relative age, their social relationship and the topics they are using, among others. In addition to variety selection, there are the skills of structuring a conversation involving such skills as who talks first, interruptions, topic changes and silences, and many others. These skills of organizing conversational interaction are referred to as **pragmatics** by linguists.

Pragmatics covers a wide range of topics, from the language choices of bilingual speakers to a parent's timing of interruptions of a child's explanation of a misdeed. Since a good part of pragmatics is as much a system of interpersonal interaction within a society (dealing with such various phenomena as male-female relationships, dominance and subordination among social superiors and inferiors, self-defense and self-aggrandizement, interactional play and manipulation—or, in other words, matters of social and cultural anthropology) as it is part of the system of producing and organizing speech signals, we will not deal with pragmatics as a separate topic. Instead, the role of speaking within a system of societal constraints will be limited to a discussion of the study of the system of choosing which speech variety to use under which social circumstances. This topic will be discussed further under the heading “sociolinguistics” below.

The Cognitive Skills for Speaking: Grammar

There are still other skills required of our two speakers. The units of meaning (“words” will do for now as a term for these units) themselves will likely be organized as particular combinations of smaller units of meaning, and words must then be combined in particular ways into groups which form sentences and phrases. These skills of forming and organizing words are called **grammar** by linguists, divided into the cognitive abilities to make words (**morphology**) and make phrases and sentences (**syntax**). You might object that these do not appear to be much different from semantic abilities. Morphology and syntax certainly have semantic results, but they are not the same skill as semantic organization.

The difference between semantics and grammar may not be readily apparent even with an example from the reader's own language because it seems natural that a particular meaning be represented by a particular word arrangement. Three kinds of experiments should convince you that speakers can combine the sound group signals (parts of words and words themselves) only in restricted ways if they are to successfully communicate—even if all the necessary semantic elements are present. First, you could try some rearrangements of signals and see if they retain their meaning effectively. (In English, for example, try “windun” for “unwind”, “lyishboy” for “boyishly”, or “man old my” for “my old man”.) Second, you could compare the way you put words together with the way you construct words to see if they are structurally and semantically consistent. (For example, in English we use phrase structures such as “the plural of book” or “more than one book” where the words indicating plurality precede the word to be pluralized but when we construct a plural word such as “books” the part of the word indicating plurality follows the

word. “Book more than one” has a different meaning and “sbook” has no effective meaning even though in both all the semantic signals are present.) A third way to demonstrate the crucial role of grammar is to attempt to translate some word or phrase in your own language into some other language. You can use a bilingual dictionary to discover all of the necessary units whose combined presence should express an effective translation, but unless you know how to organize these units in the appropriate manner the dictionary information is not too useful.

A person wishing to translate the English phrase “He is hitting him” into Kiswahili might locate the following Kiswahili morphemes: “he”= a, “is...ing (present tense)”= na, “hit”= piga, and “him (third person object)”= m and thus try “a-na-piga-m.” This is close, but would not be easily understood, since the actual order in Kiswahili would be a-na-m-piga (he is him hitting).

The Complexity of Speaking

Our two ordinary adults will have to use all of these skills in order to satisfactorily conduct their conversation. **Phonology** (or the ability to make and arrange the appropriate speech sounds), **semantics** (the ability to comprehend and use the appropriate meanings), **pragmatics** (the ability to organize the interaction of speaking in an appropriate manner), and **grammar** (divided into morphology, the ability to construct words, and syntax, the ability to construct phrases and longer word sequences) are all necessary and each of these is an intricate system in its own right. Taken together they make speaking, even in this simplified overview, a wonderfully complex activity.

Yet the actual complexity of conversation is even greater than it appears. The systems which organize speaking (and other modes of conversation such as writing and signing)—taken together what we will call **language**—do not consist of a single neat collection of separate and unvarying rules. Language is best thought of as made up of multiple systems of more or less flexible guidelines which may overlap with each other, vary according to context and circumstance, and, especially, change over time.

The Study of Language and Humans

What is the interest of social science in language (studied primarily through the analysis of speaking)? Given that disciplines such as anthropology, sociology, history, and political science, among others, are the study of humanity, the answer is obvious; as a human activity,

language-directed activity must be included along with the activities associated with other social and cultural systems such as decision-making, ritual, clothing selection, tool use, classification of kinsmen. But to the field of anthropology in particular, language has more than just this general principle of inclusion to demand its study. Language is more than just one among many important human attributes; ***it is the single essential and definitional characteristic of humanity***. Without studying language, anthropology cannot claim to complete its primary mission, a holistic study of humanity.

Language and Human Evolution

When did language appear? Was it a gradual process, with language components accumulating over millennia, or did a basis for symbolic and grammatical communication appear relatively quickly, perhaps as a result of genetic change. This latter view is associated with the corresponding idea that language appeared along with the evolutionary arrival of *Homo sapiens*, within the last 1—200,000 years. This has stimulated considerable debate among linguistic anthropologists and it is a challenging issue due to the absence of any direct evidence for the presence of language in the distant past. Clearly, the presence of complex artifact manufacture that show the use of secondary implements (tools to make tools with its implied planning and conceptualization requirements) along with a set of complex artifacts would strongly indicate the presence of language. Consequently, most anthropologists would agree that by the appearance of cave decoration about 30,000 years ago language was also present. Failing the discovery of similar complexity in the artifact record at a much older age, it would seem that we are limited to only speculation about language presence earlier than this. In fact, the hundreds of thousands of years of the Early and Middle Paleolithic when tool types did not change significantly for earlier hominids such as *Homo erectus* would appear to support this idea, namely that language was not present until relatively recently. Still, it is difficult to avoid the conclusion that early hominids possessed some sort of system to manipulate concepts as part of their successful adaptive package. In addition to some sort of advanced imitative and experiential transmission of techniques to modify materials, in particular stone, the successful spread of *Homo erectus* over much of the Old World would suggest an effective use of social categories, i.e. a “cultural” tool kit. It is also difficult to imagine a cognitive system as complex as language arriving fully developed only with *Homo sapiens*, without earlier transitional systems.

Language and Primate Communication

Related to this question is the possibility of language-like communication among our closest animal relatives, the primates. Chimpanzees are arguably our closest relatives and if they can be shown to possess even a rudimentary language ability then this would argue for a very ancient presence of language. Unfortunately, research into the communicative capabilities of chimps, as well as gorillas and orangutans, has not been conclusive. It is clear that these other apes do not possess the vocal capabilities of humans. But the human ability to employ language in communication is not dependent upon sounds, as is easily demonstrated by the use of other modes of linguistically complex signaling used by humans such as signing, Braille, and writing. Researchers have been creative when working with non-human primates and, instead of employing vocalization, have used gestures (American Sign Language) or physical tokens (such as differently shaped and colored plastic pieces on a magnetic board, or something like a typewriter but with shapes instead of letters). There have been intriguing results, and some primates have acquired lexicons of two hundred “words” and apparently been able to string these together in what has been argued to be grammatical patterns. However, a human cannot have an *open* conversation with a chimp. Keep in mind that any human (one without a cognitive challenge such as aphasia or Down’s Syndrome) can have a conversation with any other human once one or the other language is learned. And this conversation can be about almost anything—family life, morality, dreams, economic activities....

The significant features of human language are that it is an “**open**” and “**productive**” system. This means that humans, as a matter of completely ordinary conversational exchanges, can create brand new messages and that these can be understood and responded to, and this can be done on any topic, even imaginary ones. All that has been learned about the origins of human language from studying the abilities of our primate kin is that this kind of ability must have developed separately along the human evolutionary line. Language wasn’t something that was carried along from our common primate ancestors. There probably was a communicational system among our primate ancestors that would allow some manipulation of limited content exchanges, but nothing more. But whatever this was, it stayed limited until at some point in human evolution a fully creative capability arose that permitted our kind of primate to produce essentially unlimited content exchanges.

Linguistics within Anthropology

One of the four major subdivisions of anthropology, therefore, is linguistics. It touches on many other areas that anthropology is interested in such as primatology, human paleontology, and the archaeology of ancient humanity. Language is such a seminal subject, having an effect on all human endeavors, that many other disciplines besides anthropology include it as part of their field of study. At a growing number of universities it exists as a separate discipline. This book reflects the author's anthropological background and training and thus presents linguistics from an anthropological perspective. Not, however, from the viewpoint of the specialized fields labeled as *anthropological linguistics* or *linguistic anthropology* (which deal in a variety of ways with the manifold and fascinating uses of language by humans or its conceptualizations by humans to a large variety of ideological ends), but as an introduction to the study of language structures hereafter just called linguistics. Simply stated, linguistics is the study of language²

Given the importance of a broadly comparative approach, one that applies to all humanity, linguistics gathers and studies data on as many languages as possible, used by any community for any communicative purpose, but especially the ordinary speech interactions among a community of speakers. While it is interested in systems of writing, it does not limit itself just to those languages associated with written traditions, nor just to those languages associated with economic and political power. And although there are many applications of linguistics, from designing writing systems to speech therapy, and while several of these applications are also of particular interest to other fields, such as historical reconstruction, speech therapy, and cross-cultural translation, linguistics itself is primarily interested in describing and explaining the operation of language in human communication rather than the applications of this knowledge, however practical they may be.

Subdivisions of Linguistics

Linguistics may be divided into a number of mutually dependent areas of study. These are examining the organizational systems of one language,

² The term "language" is often used very broadly to refer to any form of communication, e.g. computer language, bee language, dolphin language, the language of the flowers. While communication is present throughout the living world (indeed one definition of "living matter" is the ability to communicate), the term "language" should be reserved for one particular, very complex, form of communication—that which organizes human communication.

or what we will refer to as **descriptive linguistics**, trying to figure out how language systems change over long time periods and reconstruct language histories, or **historical linguistics**, figuring out the processes by which human acquire language, especially in the first several years of life, or **developmental linguistics**, and the manner in which human communities value and allocate their language behavior or **sociolinguistics**. Linguistics itself can also be viewed as being part of a general study of all animal communication, in particular the study of primate communication. But however fascinating it would be to compare animal communication systems, even just primate systems, we will not go into this topic beyond what was mentioned above. The position taken by the author is that human language-based communication is such a distinctively complex system that little is gained by comparing it to other systems in an introductory text. This topic has seized the imagination, and unfortunately the anthropomorphizing inclinations, of so many people that the communication of chimpanzees sometimes receives as much attention in introductory anthropology texts as does all of the diversity of human systems. Without denying the suitability of this type of comparative study, it will be a sufficiently challenging task here to present an overview of language.

Descriptive Linguistics

Descriptive linguistics is the description of the language systems in use for a particular community of speakers within a particular time frame (a generation or two). As you should appreciate by now, descriptive linguists face a challenge worthy of the most diligent, ingenious and resourceful investigators. In fact, there exists no complete description of any community's language, and the hundreds of partial descriptions are the work of many researchers, each contributing analyses on different aspects of a language system. The field of descriptive linguistics itself is further divided into such areas as phonology, grammar, semantics, and pragmatics. The comparative mission of linguistics would not be possible without the data provided by descriptive studies; those who use this material to draw conclusions about the state of human language can even be considered to be in another subfield—the study of language universals. One significant conclusion reached from this approach is that the approximately 1000 languages which have been studied in depth (of the 7000 or so languages in existence during the last several centuries) are all of equivalent complexity, that is to say that their organizations are equally complex and that each of these provides a complete basis for all of their respective community's communicational requirements.

There are a fairly wide range of descriptive approaches. Language, in what should become apparent over the course of this text, is an exceedingly complex phenomenon. As such, there are many different, but somewhat complementary, approaches to describing it. In the early part of the last century, there was much interest in system of sounds, but by the end of the century syntax and semantics had come to the foreground as the proper focus of linguistic study. Obviously, arguments over how best to describe languages veer over into theoretical arguments about what constitutes the general state of linguistic systems in the mind. We shall not consider these theoretical arguments here but hope that the reader will become interested enough to pursue these questions further.

Historical Linguistics

Historical linguistics studies the changes in language systems over time, as measured in generations. This is to differentiate it from developmental linguistics which studies the changes in language abilities through the human life cycle, particularly during childhood when language is substantially acquired.

Historical linguistics is not directly concerned with language change as it occurs. Rather, this branch of linguistics is interested in reconstructing the changes that must have occurred in the past and which now can be deduced by comparing selected contemporary languages. Using the data provided by descriptive linguistics, historical linguists look for evidence that two languages must have shared the same distant language “parent.” Languages are such complex systems, that if two languages share common features (and if extensive borrowing by one from the other can be ruled out) then this is assumed to be the result of their having developed from the same language rather than the result of chance. One result of this investigation is to group languages into “families” and “stocks” according to common parentage. For example English and German are members of the Germanic language family while French and Spanish are members of the Romance (or Italic) family, and both, together with many other groupings such as Balto-Slavic (which includes Polish and Russian) and Indo-Iranian (which includes Persian, or Farsi, and Hindi) are part of the Indo-European language stock.

Historical linguists also try to place this classification into a time frame by determining when languages diverged. Some would be closely related “sister” languages which diverged fairly recently from the same “parent” form, and some would be distantly related “cousins”, having “parent” forms which themselves diverged long before the others—in other words finding language “genealogies.” In addition, by comparing vocabularies of related languages and thereby reconstituting the vocabularies of parental

languages, historical linguists attempt to determine aspects of the cultures and culture history of the speakers of the past languages. The reader should appreciate how fruitful this approach to language can be to other fields of anthropology, especially archaeology.

Developmental Linguistics

Developmental linguistics, on the other hand, examines language through the life cycle of individuals, in particular during childhood. One of the most remarkable aspects of language is its acquisition by children—who, within the first few years of life, as immature and cognitively unformed persons, and with little or no specific training, acquire this intricate set of cognitive systems for open and productive communication. Of course, as animals, we are specialized to learn language. Developmental linguists try to determine the exact nature of how language is acquired—the order in which different language systems are acquired, the role of the child’s own creativity and the influence of different primary language user models, especially the differing effects of peers and adults and the differing effects of user models who speak different speech varieties, among many other topics.

Sociolinguistics

Finally, perhaps the most comprehensive of the sub-fields is sociolinguistics. Speaking itself is organized into types of **speech varieties** and these are interrelated with kinds of speakers and hearers, kinds of social situation and kinds of topics. Furthermore, different ways of speaking carry different evaluations, from honor to disgust. Language systems can be codified into written standards whose use can even be legislated and enforced. Communities can deliberately attempt to change the way they use speech varieties (especially when they are based on different languages) as part of developing their own symbolic image. Ways of speaking can atrophy and die out, or grow in importance from being merely a “dialect” to becoming the “Language” of a great empire. Sociolinguistics attempts to study all of these ways that speaking and language, as social institutions and objects of belief, are the focus of social activity. It deals with the dynamic existence of multiple ways of speaking within individuals (bilingualism) and societies (linguistic pluralism). More than any of the other sub-fields, sociolinguistics overlaps with other social science disciplines such as sociology, psychology and political science.

Terms to Know for Chapter One

phonetics	syntax
articulation	language
phonology	linguistics
semantics	openness and productivity
pragmatics	historical linguistics
speech varieties	developmental linguistics
grammar	sociolinguistics
morphology	descriptive linguistics

CHAPTER TWO

PHONETICS: THE DESCRIPTION OF SPEECH

The study of language must have a firm base in the accurate description of speech. This chapter will introduce articulatory phonetics as a method of describing speech based on the positioning and operation of the various organs in the vocal tract. Some of these organs are familiar, such as the tip of the tongue, some you may not have heard of, such as the uvula, and most you know of, but by a different name, such as the palate and velum which are the roof of the mouth. Following a brief discussion of the acoustic properties of speech, we will plunge into learning the many parts of the vocal tract and the ways in which they act to produce different kinds of speech sounds. We will then learn a system of phonetic notation which can serve as a convenient means of recording speech sounds according to their articulatory characteristics.

The Importance of Speech

Linguistics, as a social science, must have a means of accurately and reliably describing the phenomena it studies. The subject matter of linguistics is language, which involves a set of mental systems that organize speaking (and hearing). The state of cognitive investigation is not yet capable of describing all the particular language activities occurring as neurological processes within the brain. Linguistics, therefore, must do the next best thing and that is to carefully describe speech as the physical manifestations of language and then to assume that the patterns which can be determined in observable speech events actually reflect the mental processes of language. This is an important assumption and cannot be emphasized too strongly. The linguistic anthropologist, especially, is committed to an approach in which statements about language have to be based in valid descriptions of speech. This means that linguistics must include a reliable procedure for recording, transcribing, and analyzing speech sounds. This is called **phonetics**.

Acoustic Phonetics

Speech as Sound

Speaking is a physiological activity producing audible sounds. There are two ways of describing speech. Speech sounds can be described using their acoustic, or physical, characteristics or by using the way these are articulated in the vocal tract. The latter method is commonly used by linguists in the field. There are several reasons for this, but before defending an articulatory description let us briefly examine an acoustic description. Speech is sound and sound consists of vibrations in some medium; in the case of speech, this is typically air.

An acoustic description is consequently a precise account of the nature of these vibrations: how rapid they are (“frequency” expressed in complete vibrations—i.e. the air particles pushed to one side, springing back to the other side and then coming back to their original position—per second), how much energy they have (loudness, or the “amplitude”—i.e. size—of the vibrations, expressed in decibels), how long they last (time, expressed in seconds or, more usually, fractions of seconds), and how complex the vibrations are. This last feature refers to the fact that speech sounds rarely, if ever, consist of pure oscillations such as one might get by striking a tuning fork whereby the air is pushed smoothly back and forth in its vibrations. The vibration patterns in speech are a result of the airstream’s many turns, constrictions, and bumps as it moves through the vocal tract. Furthermore, the several regions in the vocal tract act to enhance some vibrations and dampen others. As would be expected, by the time the airstream emerges it carries the effects of many vibrations. Acoustic descriptions break down the complex totality of a speech sound into its harmonic constituents and also specify the energy level of each of these. An instrument called a spectrograph (see Figure 2–1) produces a visual printout of segments of speech (about 4 seconds). This kind of information can also be represented (and stored) electronically using the mathematical properties of the vibration.

In the two examples in Figure 2–1, speech is represented by a large number of vertical bands. Each band is the result of a series of tiny light bulbs in a vertical line over which light sensitive paper is passed. Each bulb will respond according to speech sound energy in a certain frequency range, and the intensity of the light given off by the bulb is proportionate to the amount of energy in that frequency range. A darker part of the vertical band indicates greater energy at that frequency range, a lighter or white area indicates less energy or the absence of energy at that frequency range. Each type of sound has a distinctive pattern of dark bands over its brief time of duration. The two samples of speech upon which the

spectrograms are based are represented by normal spelling above each one, however the spacing has been adjusted so that the letters representing sounds are above the corresponding bands. (adapted from Denes and Pinson)

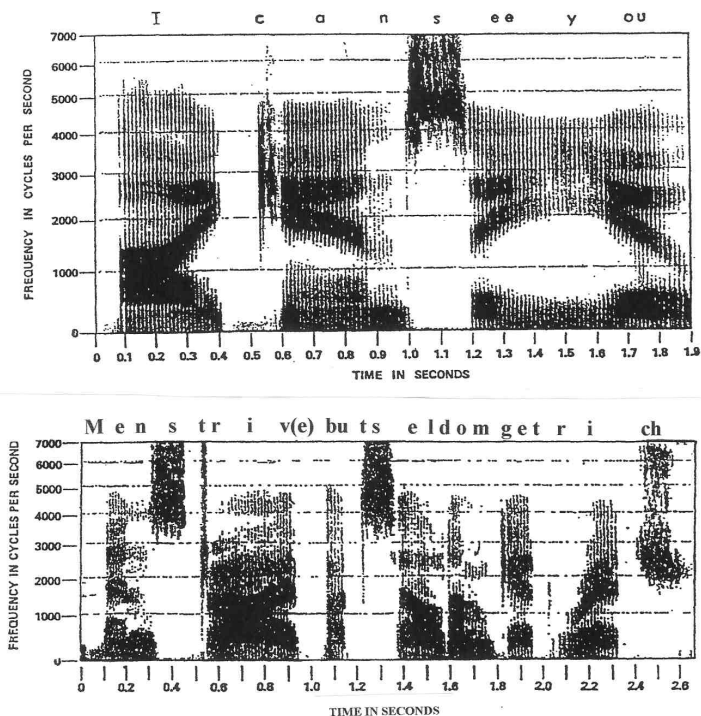


Figure 2-1 Spectrogram of Human Speech

This is what acoustic scientists use when they study the transmission possibilities for speech (e.g. radio, tape recording, telephone). Needless to say, these acoustic descriptions of speech, employing electronic sensing equipment and mathematical notation, appear to have the advantage of great precision and objectivity. It will be made much clearer below, but at this point you can assume that an articulatory description, using the positions, movement and shape of the various organs of the vocal tract, is not as precise as an acoustic description. So why do linguists use articulation as the basis for their descriptions of speech?

Laboratory Phonetics

A related development in the study of human speech is the use of sensing equipment that are capable of precisely measuring articulatory activity by the vocal tract organs. One example of this approach is the use of a device that is inserted in a speaker's mouth and covers the roof of the mouth (palate). This "pseudo-palate" has a large number of contact electrodes on its surface and can therefore precisely measure where the tongue is making contact with the palate. Other methods use x-rays to view the precise positioning of parts of the vocal tract. Laboratory phonetics shares with acoustic phonetics the device-assisted refinements that provide for precise descriptions of speech.

Advantages of Articulatory Phonetics

Precision in recording the vibrations of speech sounds and describing their exact placement in the vocal tract is valuable for a number of reasons, including acoustic analysis, speech synthesis, and diagnosing speech disabilities. But it takes special training to interpret these results and, therefore, does not lend itself to the learning and acquisition of speech sounds directly. For the human speaker, individual sounds are identified not so much by their absolute physical properties as by their relative differences from one another.

Besides, an articulatory description is more "human." Humans do not learn to speak by consciously manipulating frequency, harmonics, and energy, but by manipulating their vocal organs to produce speech sound differences. In other words, much as the acoustic characteristics of speech are important, the primary concern of speaking is the activity of articulation.

The application of acoustic phonetics was once constrained by its reliance on expensive equipment that was delicate and bulky. One virtually needed a roomful of machinery to conduct acoustic analyses just a few decades ago. Back then, field researchers would very much like an acoustic analysis to reveal the intricacies of some speech sounds from time to time so that they could better decide how to move on with their analyses. But until several years ago, this was a luxury that one could only dream of. Today, an acoustic analysis is very much at a field researcher's finger tips if he or she is equipped with a digital recorder or camcorder in addition to a laptop installed with acoustic software. Laboratory phonetics, by its very name, must be done in a laboratory setting using often bulky machinery.

For our purposes, acoustic and laboratory phonetics are best viewed as valuable ways to describe speech activity that are complements to, rather

than replacements for, articulatory phonetics. These descriptions are meant for a trained eye, in pursuit of technical questions involving speech production, not the introductory linguistics student. If acoustic and laboratory phonetics present views of speech sounds in scientific terms, articulatory phonetics presents an empirical approach to the description of speech sounds in behavioral terms.

Articulatory Phonetics

Articulation

An articulatory description consists of describing speech sounds according to the configurations and movements of the many parts of the vocal tract. Since it is tied to physiological landmarks (e.g. lips, tongue, nasal cavity) the speech of any human can be described using the same system since all humans have the same landmarks and potential configurations of their vocal tract. An important corollary of this fact is that it is possible for any human to speak any language—there are no special physiological requirements for any language. A descriptive system for articulation should be so designed that it is applicable to all human speech behaviors. If, for example, we were to omit the possible configuration of the root of the tongue against the back of the throat (the pharynx) simply because this configuration does not occur in English then this would severely limit the usefulness of our system. If we attempted to describe a language which has this configuration for some of its sounds (e.g. Arabic) our English-based system would not be sufficient and would have to be modified if it were to be useful. More importantly, the initial absence of this configuration in the phonetic system, if it were being applied to Arabic, might bias the linguist from accurately describing this kind of configuration, and many others not present in English, especially at the beginning of the study.

Manner of Articulation

Let us begin with the general ways in which a stream of air used in speech can be modified by the organs that make up the vocal tract. We will then examine the various places within the vocal tract where these modifications occur. Following this we will return to a more detailed discussion of manners of articulation. A note of caution is necessary here, however. We will only examine a “bare bones” version since the purpose of this book is only to present an introduction to the description of human articulation. The reader interested in pursuing phonetics further should start by carefully examining the International Phonetic Alphabet given in

Appendix C, and the bibliographic sources listed (particularly Abercrombie, and Samarin).

Speech is generated with an **airstream** coming into or out of the lungs. This airstream is forced out or sucked into the lungs by the operation of surrounding muscles. Modifications can be made in either direction of the airstream, outward from the lungs (pulmonic egressive) or inward to the lungs (pulmonic ingressive). Ingressive speech sounds are present in some languages (notably the Khoisan languages of southern Africa), but unless otherwise indicated we will only be describing egressive sounds since these are much more typical of human speech. We can examine modifications of the airstream by considering what is done to it in the vocal act. One type is a complete interruption of the air flow, another is simply a shaping of it. In between there are many different kinds of modifications.

A complete blockage of the airstream gives rise to a **stop**. In many cases it is the nature of the release as much as the blockage itself that characterizes the speech sound. This is why some linguists use the term plosive—as in “explosion”—since it emphasizes a sudden release of the airstream. Stops, then, are made by a closing and reopening of some part of the vocal tract. This can be done at a number of points in the vocal tract. The closure typically occurs in a single place but it may occur in more than one place (e.g. Yoruba spoken in West Africa has many words that have stop sounds with double closures).

Another manner of articulation produces a **fricative**. In this case, the airstream is not stopped, but squeezed through a narrowed opening in the vocal tract. This creates a considerable amount of turbulence, which is an important acoustic feature of these sounds. You might remember the term fricative by thinking of “friction”. (These sounds have also been referred to as “spirants” and “sibilants” in some phonetic systems.) Fricatives can be made all along the vocal tract.

Another manner of articulation produces a **resonant**. For this sound the airstream is not squeezed but instead directed around some part of the vocal tract. This is done in three different ways. The tongue can be bunched up in the mouth with the airstream channeled over it, the front of the tongue can be raised up against the roof of the mouth and the airstream channeled around both sides, and the mouth can be closed off with the airstream directed through the nasal cavity. Stops, fricatives, and resonants, taken together as a group, are called **consonants**.

The last manner of articulation produces **vowels**. Here the airstream passes through the mouth and is shaped by varying positions of the tongue. These are also subject to the configuration of the lips and whether or not the airstream is also passing through the nasal cavity. Vowels will be described below.

At this point readers might like to try and use what they have learned and identify which kinds of consonantal manners are used in their own speech. Problem 1 in the exercises for this chapter (hereafter given by chapter and exercise number: 2–1) asks you to say your name and then list the manners of articulation for the consonants that occur (omitting, for now, vowels). Read the guidelines and pay attention to the examples provided from the author’s own speech.

The Vocal Tract: Articulatory Landmarks

As can be seen from the above, there are many ways of modifying the airstream as it passes through the vocal tract. This section will deal with **places of articulation**, where these modifications occur. The vocal tract is not made up of separate segments, easily distinguishable from one another. What we shall do is to point out the areas and organs which are commonly used as landmarks in describing the articulation of speech sounds. Figure 2–1 shows a cross section through the vocal tract. The airstream exits the vocal tract at two openings: the lips and the nostrils. Unlike the nostrils, the upper and lower lips are active articulators and serve as landmarks. When one or both lips are used in articulation the adjective **labial** is applied to the sound. Behind the lips are the teeth, the next landmark. When “teeth” (usually the front teeth) are described in articulation with the adjectives **labiodental** and **dental** are used.

The roof of the mouth extends from just behind the upper front teeth to where the throat branches into the nasal and oral cavities. This region is divided into several places of articulation. From behind the teeth to about midway back, there is a bony backing to the roof of the mouth (shaded in Figure 2–1) and this is correspondingly called the “hard” palate. The very front of the hard palate—the gum ridge bulge just behind the front teeth—is the **alveolar ridge**. When this is involved in articulation the adjective **alveolar** is used. The rest of the hard palate is just called the **palate** and when it is involved in articulation the adjective **palatal** is used. Some articulation positions are slightly behind the alveolar ridge, and overlap with the front part of the palate. These are termed **alveo-palatal**. That part of the roof of the mouth without the bony backing is called the **velum**. At the very end of the roof of the mouth is the **uvula**. The adjectives used for these are, respectively, **velar** and **uvular**. The very back of the vocal tract, in the throat, contains the larynx at the top of which are the **vocal cords** or glottis. Stop and fricative articulations may be produced here and for these the term **glottal** is used.

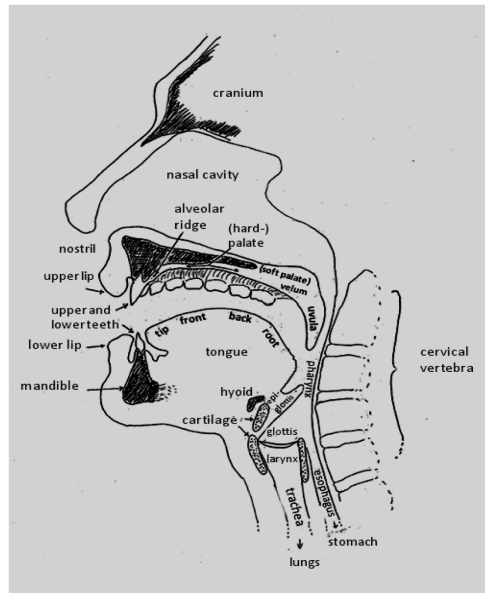


Figure 2–2 Cross-Section of Human Vocal Tract

The bottom of the mouth itself is not important in articulation, serving only as the place for the muscle attachments of the tongue. It is the tongue which is the dominant articulator of the mouth, although not for the entire vocal tract. (The importance of the tongue is suggested by the use of “tongue” as a synonym for language.)

The tongue can take a variety of configurations by itself and together with the lower jaw’s (mandible) capability to be raised and lowered it can greatly affect the nature of the airstream through the mouth. The tongue is divided somewhat arbitrarily into convenient landmarks (in contrast to the more ‘natural’ divisions of the roof of the mouth). The front end of the tongue is called the **tip** (or the apex by some linguists). The tip includes the entire conical surface, top and bottom. The rest of the tongue’s landmarks refer to only the top surface. These divisions vary among different phonetic traditions, compounded by there being no obvious physiological boundaries, and so we will present a basic series of landmarks. The part behind the tip is called the **front**. When the tongue is at rest its position would roughly correspond to the hard palate, with the tip located behind the front teeth. Behind the front is the **back**, which would correspond to the soft palate, and finally the **root** of the tongue which is opposite the back wall of the throat. The throat, **pharynx**, lies to the rear of the mouth, or oral cavity. It extends from a little above the top

of the velum and uvula, where it goes into the nasal cavity, down to the top of the vocal cords in the **larynx**. The vocal tract extends to, and includes, the larynx, and it is connected through the windpipe (trachea) into the lungs which provide the driving force of speech in the airstream.¹ Now that we have the landmarks, we can proceed to a discussion of how these are used to describe the articulation configurations for consonants and vowels.

Places of Articulation: Consonants

A convenient way to describe consonant articulations is to think of them as the product of two articulators affecting the airstream—a movable articulator that is in contact with, or very close to, an immovable articulator. The movable articulators are the lower lip and the tongue, and the immovable articulators include the upper lip, upper front teeth, and the roof of the mouth.²

Table 2–1 Descriptive Terms for Places of Consonant Articulation

Term	Movable Articulator	*	Immovable Articulator	English Examples
bilabial	lower lip	→	upper lip	“ <u>b</u> ut” “ <u>h</u> im”
labiodental	lower lip	→	upper front teeth	“ <u>f</u> oot” “ <u>p</u> ave”
dental	tip (of tongue)	→	upper and lower front teeth	“ <u>t</u> he” “ <u>p</u> ath”
alveolar	tip (of tongue)	→	alveolar ridge	“ <u>d</u> rink” “ <u>b</u> us”
alveo-palatal	front (of tongue)	→	alveolar ridge and palate	“ <u>sh</u> oe” “ <u>ca</u> ch”
palatal	front (of tongue)	→	palate	“ <u>k</u> it” “ <u>t</u> ree”
velar	back (of tongue)	→	velum	“ <u>g</u> um” “ <u>r</u> ing”

¹ The esophagus, or the tube leading to the stomach, also enters the pharynx (next to the larynx), but except for some very special role in sound production such as burps—which are produced for a variety of purposes, including being a source for voicing for those who have had their larynx removed—the esophagus is not included in the description of speech articulation.

² There are exceptions to a neat system of movable and immovable articulators: the upper lip is not “immovable” in the same sense that the roof of the mouth is immovable. Its ability, along with the lower lip, to pucker, for example, is an important feature of vowel articulation. Furthermore, a good part of the “movement” of the lower lip and the tongue toward and away from an immovable articulator is in fact movement of the lower jaw (mandible), yet the mandible is not considered a part of the vocal tract. Finally, the vocal cords (the two membrane folds of the glottis—the opening in the larynx) are both movable and move toward and away from each other. Still, this system is reasonably useful as a general means of categorizing places of articulation.

uvular	back (of tongue)	→	uvula	“cough”
glottal (laryngeal)	both membrane sides (vocal cords) of the glottis (the opening in the larynx)			“ <u>h</u> it”

*The arrow is to be read as: The movable articulator is close to or touching the immovable articulator (the Place of Articulation).

Table 2–1 presents eight commonly used articulatory configurations. Each of these is given a descriptive name, usually based on the name of the immovable articulator, or **place of articulation**. Each row lists the configuration name, movable articulator involved, and immovable place of articulation. Examples of these from English words are also supplied.

Now readers might like to try and use what they have learned and identify the places of articulation that are used for the consonants in their name. Exercise 2–2 asks you to say your name again as for 2–1 and then for each consonant list where in the vocal tract (place of articulation) it is made (again, omitting vowels). Reread the guidelines and pay attention to the examples provided from the author’s own speech.

Basic Consonant Chart

When the basic manners of articulation are combined with the different places of articulation they produce an articulatory chart on which any consonant sound can be located. Figure 2–3 shows such a simplified consonant chart. The places of articulation are listed across the top, with the bilabial position on the left and then proceeding to the right (as if we are going into the vocal tract of a speaker facing to the left) until the glottal position is reached. Manner of articulation is listed on the side with stops at the top and resonants on the bottom and therefore, the side listing is from greater obstruction down to lesser obstruction.

Now let’s apply both manner and place of articulation and identify the consonants in some common words—your choice. Exercise 2–1 in the Exercises for Chapter 2 asks you to say your name and then list the manners and places of articulation for each of the consonants.

		ARTICULATORY CONFIGURATION								
MANNER OF ARTICULATION		BILABIAL	LABIODENTAL	DENTAL	ALVEOLAR	ALVEO-PALATAL	PALATAL	VELAR	UVULAR	GLOTTAL
	STOPS									
	FRICATIVES									
	RESONANTS									

Figure 2–3 Simplified Phonetic Chart for Consonants Showing Place and Manner of Articulation

(Note that this chart only illustrates the articulatory dimensions of consonants, thus each box, for now, is empty)

Places of Vowel Articulation

For vowel configurations let us look again at the cross-section of the vocal tract and pay special attention to the relative position of the tongue in the oral cavity. Vowels are categorized by tongue height and by the location of the highest part of the tongue. The different positions of the tongue affect the size and shape of the oral cavity and thus the nature of the vowel sound.

A **high** position is when the tongue is raised so that the oral cavity above it is small. A **mid** position is where the tongue is more flattened so that the oral cavity above it is larger. And a **low** position has the tongue flattened more so that the oral cavity is at its most open. (It should be noted that these tongue movements are usually associated with some corresponding raising and lowering of the lower jaw; no speech sound requires that the mandible be lowered and the tongue be flattened as far as possible—in other words that the mouth be wide open.)

In addition to relative tongue height, the relative position of the highest part of the tongue's upper surface is also used. The tongue can be shaped so that the leading part of the tongue, the middle part, **or** the back part is highest in the mouth. A **front** position is where the leading part of the

tongue (approximately the tip and part of the front) is highest, **central** position where the middle part—the rest of the front up to the back—is the highest, and **back** where the back part of the tongue is the highest. (Note that this isn't quite the same as our previous division of the tongue for consonants into tip, front, back, and root.) The effect of these different tongue contours is to divide the oral cavity into two differently sized cavities. A front position means that the whole oral cavity is divided such that the cavity in front of the raised part of the tongue is smaller than the cavity behind it. A central position has roughly equal cavities in front and behind, and a back position has the smaller cavity behind and the larger cavity in front. The different resonance effect of the two cavities, in part, produces the acoustic characteristics of different vowels.

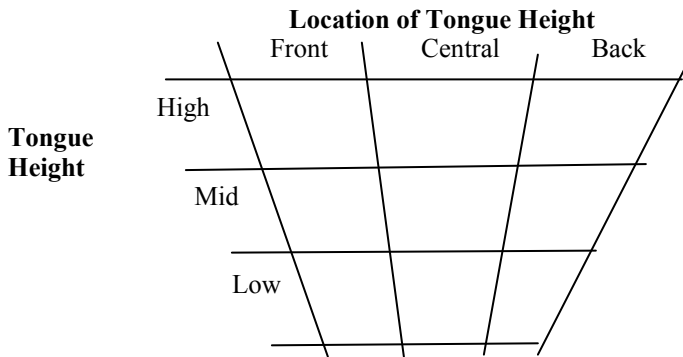


Figure 2-4 Simplified Phonetic Chart for Vowels Showing Tongue Height and Tongue Height Location

(Note that this chart only illustrates the articulatory dimensions of vowels, thus each box, for now, is empty)

Basic Vowel Chart

Combining tongue height and contour we get nine configurations. As shown in Figure 2-4, these are best represented physiologically not as a rectangle but more like a trapezoid due to the fact that there is more room from the front to the back positions at the high tongue position than at the low position.

Now let's try to identify the vowels using the simplified vowel chart (vowels are more difficult than consonants, and as you will see below there are several more features that will be needed, but do the best you can). Exercise 2-2 asks you to say your name as in 2-1 and then list the kinds of vowels that occur (omitting the consonants). Reread the

guidelines and pay attention to the examples provided from the author's own speech.

Now let's apply both manner and place of articulation for consonants **and** position for vowels in some common words—your choice. Exercise 2–3 asks you to say the names of four things in your room and then list the manners and places of articulation for each of the consonants, and position for each of the vowels.

Expanded Phonetic Charts: Consonants

You should now have the basic descriptive outlines of consonants and vowels. However, there are several refinements that must be presented—in fact, you should have found your attempts in the exercises a little challenging because some aspects of your speech did not seem to fit neatly into the simplified systems so far presented, especially for the vowels.

We will start with consonants. Figure 2–3 presented just the basic outlines of a consonant chart. In order to make our descriptive system more broadly useful we will need to expand the chart.

One additional feature that is needed is called **voicing**. Consonants frequently (but not necessarily) include an important simultaneous manner of articulation. This is rapid vibration of the airstream produced by motion of the vocal cords in the larynx (by itself, this vibration sounds something like what you do when your doctor asks you to open your mouth and say “aahhh”). If the muscles of the vocal cords are tightened just enough, the air pressure in the lungs will have to build up a little to push them open for a small part of the airstream to pass through. Immediately afterwards (since the air pressure has now slightly decreased), they ‘snap’ closed only to be pushed apart again. When this process is done a hundred times a second³ it produces a tone somewhat like the vibrating reed(s) of a musical instrument or the vibrating arms of a tuning fork. This is considered a secondary manner of articulation that affects stops and fricatives, and so each manner of speech sound on our simplified chart (that is, each row) will be divided into a **voiced** (vd) column, if the vocal cords are tense so that they are vibrating, or a **voiceless** (vl) column if they are loose so that the airstream passes between them without marked vibration. Instead of just stops and fricatives, for example, we now have voiced and voiceless stops⁴ and voiced and voiceless fricatives—by convention, the voiceless

³ Humans have the capability to greatly alter the frequency of this vibration but this variation is not necessary for describing the feature of “voicing.” It will however, be discussed below under “tone.”

⁴ Actually there is a crucial difference between stops and fricatives in regard to voicing. Fricatives can be voiced or voiceless during their entire production. I.e. a

column is on the left side and the voiced on the right. In keeping with presenting only a basic descriptive system, voicing is only indicated for stops and fricative, not resonants (which, along with vowels, are usually voiced).

Two additional features will be added to the stop row. The first is **aspiration**. This refers to how forcefully the airstream is pushed at the release of the vocal tract closure. When there is extra pressure on the stopped airstream from increased muscle tension the airstream is pushed strongly out when the closure is opened. This produces an audible effect because this momentary rush of air contributes to the airstream's vibrations. (You can also feel aspiration by placing your hand just in front of your lips while you say the nursery rhyme "Peter Piper picked a peck of pickled peppers..."—paying heed to the beginning of each word starting with a "p".) Unlike voicing, aspiration can be graded according to its force (no, weak, and strong aspiration). The approach we will use here to add aspiration to the chart will be to divide the stop row into two, the top for **unaspirated stops** (sometimes referred to as "simple" stops) and below this a row for strongly **aspirated stops**.⁵

The second modification to stops is called **affrication**. This produces a sound which seems to be a combination of stop and fricative. These are stops which are not released "cleanly", but instead have a fricative-like release. In other words, at, or slightly behind the point of closure, the vocal tract organs separate to release the closure but slowly enough so that as the airstream resumes its movement along the vocal tract it is squeezed through a narrow opening. In some cases, especially where the constricted release is at a different place than the stop, this can be treated as two separate consecutive manners of articulation—a stop followed by a fricative. However where both the stop and the constricted release occur at the same place this is more usefully treated as a single sound, a stop

voiced fricative has the vocal cords vibrating as soon as the constriction occurs and this lasts until the constriction ends and the fricative is over. A stop, on the other hand, by its very nature (complete stoppage of the airstream) can only be voiced or voiceless upon its release. Technically, therefore, a voiced stop is really a stop with a momentary voiced release—i.e. when the closure is opened and at the moment that the airstream resumes is when the vocal cords are vibrating or not.

⁵ Aspiration is not easily worked into our phonetic chart and the decision to place aspirated stops in a separate row was made to give aspiration a definite "place" on the chart. However, this does not mean that it cannot also occur as a co-articulation with other kinds of stops, especially affricates. Personally, I think the best solution physiologically would be to subdivide the chart further so that each voiced and voiceless row include aspirated and unaspirated rows (perhaps omitting weakly aspirated). But this has the great disadvantage of making for a very cluttered chart. Both voicing and aspiration may co-occur in the same sound.

having a particular type of release. This is called an affricated stop, or **affricate**. These will be given the third row under stops. (Affricates may also be aspirated—see footnote 5)

Some other co-articulations can be mentioned in passing at this point, but these will not be added to the expanded chart. A labialized co-articulation refers to the presence of bilabial tightening or closure while another articulation is occurring in the oral cavity. Thus a velar stop could be labialized if it also exhibited bilabial “puckering” (the term “rounding” has been traditionally reserved for vowel labial co-articulations). The reader should be able to feel the difference between a nonlabialized velar stop and a labialized velar stop by saying “kill” and “quill” and sensing for the latter word the puckering of the lips. A glottalized co-articulation refers to a co-occurring closure in the glottal folds of the larynx. In co-occurring closures, one closure—the primary one—opens first followed quickly by the co-articulating closure. Glottalized stops are fairly common in many languages, usually with the glottal closure following the primary closure above it in the vocal tract, but this, as well as the other co-articulations, need only be briefly touched upon in this introductory text. These co-articulations will not be assigned a separate row in the expanded chart (unlike what was done for stop aspiration). Instead, as is done for aspiration in fricatives and resonants, the feature will be marked by a written device—see below under “diacritic” marks for phonetic notation.

Finally, the resonant row will be divided into three different types. Recall that resonants are created by a redirection of the airstream rather than a stoppage or squeezing of it. A **nasal resonant** is produced when the oral chamber is completely closed off, as if a stop was being made, but with the uvula depressed so that the airstream is channeled completely through the nasal cavity. The combined acoustic effects of the shape of the closed oral chamber (which can be a larger resonance area with a bilabial closure or a smaller one with a velar closure) and the passage of the airstream through the nasal chamber produce a distinctive kind of resonant.

A second kind of redirection is where the tip or front of the tongue is placed against some part of the palate and the sides of the tongue are pulled down so that the tongue is bunched into the center of the oral cavity and the airstream flows around on both sides. These are called **lateral resonants**. Finally, there are the **median resonants** in which the tongue is bunched up more than for any vocoid but not so close to a place of articulation to be a fricative. These are not easy to describe so what is given as their “place” what should be understood as an approximate location.

We are now able to present a more complete phonetic chart for consonants, Figure 2–5.

Exercise 2–4 now gives you a chance to try out your knowledge of the expanded consonant chart. Note that exercise 4 uses a different format for your description. Using one set of ten common words supplied, list the articulatory features of the consonants in the word (just writing in a V where there is a vowel). Again, reread the guidelines.

MANNER OF ARTICULATION		PLACE OF ARTICULATION																	
		BILABIAL		LABIODENTAL		DENTAL		ALVEOLAR		ALVEOPALATAL		PALATAL		VELAR		UVULAR		GLOTTAL	
		vl	vd	vl	vd	vl	vd	vl	vd	vl	vd	vl	vd	vl	vd	vl	vd	vl	
STOPS	simple																		
	aspirated																		
	affricates																		
FRICATIVES																			
RESONANTS	nasal																		
	lateral																		
	median																		

Figure 2–5 Expanded Consonant Chart

(The chart only shows the consonant dimensions so each box, for now, is empty. Shaded boxes are not possible.)

Expanded Phonetic Chart: Vowels

Several additions to the basic vowel chart must also be made. The first of these involves tongue height positioning. We started with three rows showing relative tongue height: high, mid, and low. As was mentioned above, the potential range from a low to high position in the front region of the oral cavity is greater than for the other regions and consequently it is useful to distinguish more height distinctions. Therefore, the front region is divided into two additional rows: high becomes upper high and lower high and mid becomes upper mid and lower mid. The low position remains unchanged.

The second addition involves the configuration of the lips. Both lips can be tightened in a pucker so that the bilabial opening is smaller and more circular than normal. This is termed **rounding** (this is roughly the same as “labialization” for consonants, but by convention a different term is used for vowels). One way of adding rounding to the vowel chart would be to divide each column or row in two—rounded and unrounded. The way chosen here is to have two separate charts, one for rounded and the other for unrounded. There are advantages either way but using two charts gives room for other additions and yields a less cluttered chart.

A third addition involves **nasalization**. This is where the uvula is lowered so that the airstream can go through the nasal cavity. Since the oral cavity is not closed, both the oral and nasal cavities affect the quality of these vowel sounds. Nasalized vowels are usually not given a separate place on a chart, but the author believes that it is clearer to indicate such a major modification by giving it its own space. Consequently, on both the rounded and unrounded charts, each column will be divided in two: non-nasalized and nasalized. Our vowel charts will now look like Figure 2–6.

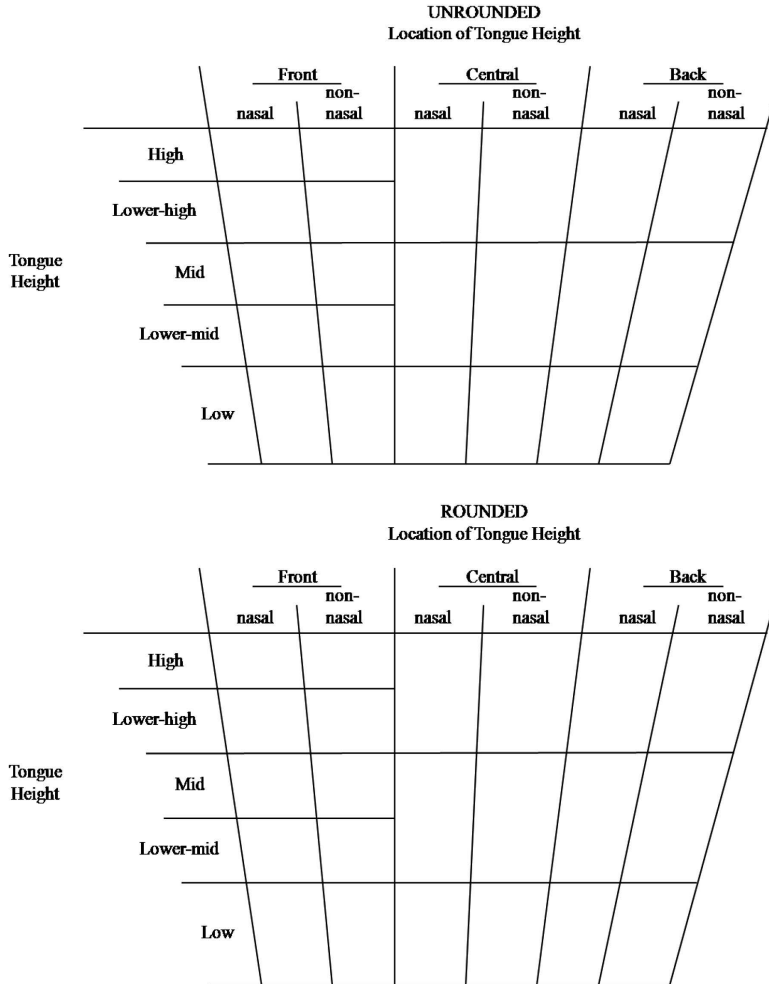


Figure 2–6 Expanded Vowel Charts

Glides

There is one final vowel attribute that must be presented. This involves a change in the configuration of the oral cavity by a tongue movement during the vowel—what is called a diphthong or **glide**. These are best presented on a separate chart using arrows to indicate the nature of the

change. One set of glides involves the tongue moving further forward and upward in the mouth from some starting position, a second with the tongue moving back and up, and a third the tongue starting from a peripheral position and moving toward a mid central position.

Glides present in U. S. English are presented in Figures 2–7 and 2–8. Figure 2–7 gives examples of what are called **forward glides**: 1 is a movement of the tongue to a further forward and upward position from a mid back position (as in “toy”); 2 is a movement forward and upward from a low central position (as in “bite”); and 3 is a glide slightly forward from a mid front position (as in “bait”). Figure 2–8 gives some examples of two other types of glides (again, these are both present in American English). 2 and 3 are **back glides**, with the tongue moving further back and upwards in the mouth: 3 is a glide back and up from a low central position (as in “bout”); 2 is a glide slightly back and up from a mid back position (as in “boat”). 1, on the other hand, shows a **central glide** (which is not present in all dialects of U. S. English) where the tongue starts at a peripheral position (such as the high front shown here) and moves to the mid central position (as in “beer” when it is pronounced like “beeuh”). Note that the 2 and 3 glides shown, respectively, in Figures 2–7 and 2–8 have the same initial positions (low central for the number 2 glide in 2–7 and low central for the number 3 glide in 2–8), thus it is not the starting vowel position which distinguishes different glides but rather what change occurs following the starting position. Glides must also be described in terms of other features such as labial co-articulation (the back glides, such as 2 and 3 in Figure 2–8, in American English co-occur with increased rounding and relative muscle tension and airstream force (stress) from the initial to the ending glide position (the central glide 1 in Figure 2–8 begins with more muscle tension and stress and proceeds to a more relaxed, low stressed condition).⁶

⁶ Glides are often interpreted as simply two vowels in sequence with the “gliding” being only the normal movement from one articulatory position to another. Sometimes they are interpreted as a vowel plus a semi-vowel (median resonant) consonant ending. Perhaps the best way for our purposes here is to think of them as “complex vowels” in contrast to the other “simple” vowels which have a single articulatory configuration.

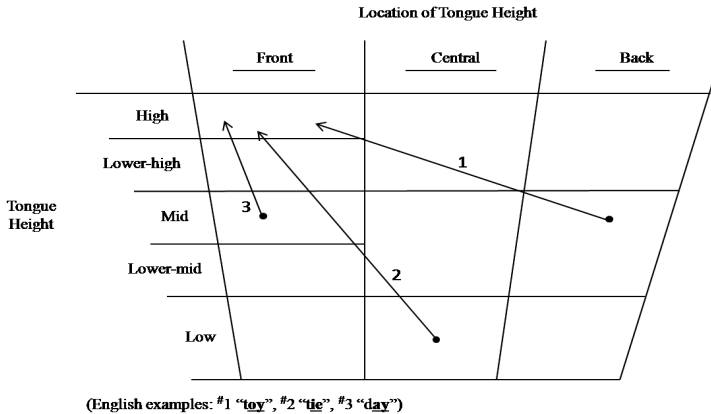


Figure 2-7 Forward Glides

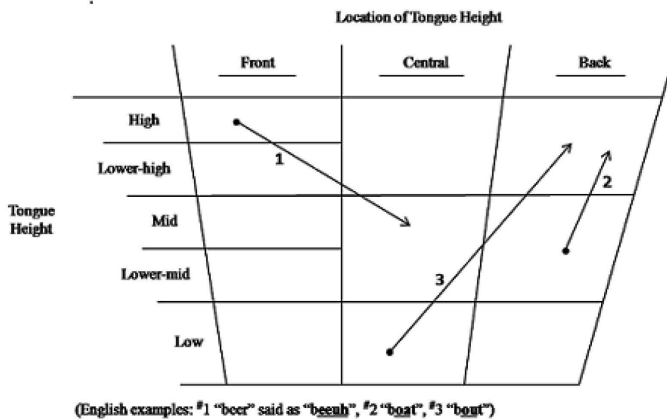


Figure 2-8 Mid and Central Glides

Vowels may also vary according to how long they are held relative to other vowels. This characteristic is called length. It may also be applied to consonants, especially non-stops. It is not a significant vowel characteristic in English (although it is not uncommon to find people mistakenly referring to short and long English vowels—a confusion linked to spelling with one as opposed to two letters). Here we will only mention it and not place it into our expanded vowel charts.

Now you may have more confidence at describing vowels. Exercise 2-5 has you using the same set of ten words for exercise 2-4 but now

describing all the articulatory feature of the vowels. Again, reread the Guidelines and refer to the examples given at the beginning of the exercises.

Articulatory Description of Selected English Words

Now that we have the basic terms for articulatory landmarks and configurations let us examine a number of examples from the author's speech where these terms are applied. Each word has one consonant or vowel part in bold and underlined and an articulatory description provided for each of these.

This book is not written for the purposes of having readers actually mastering phonetic transcription. The author, however, believes that phonetics is best presented with some practice of trying to make different kinds of speech sounds and thereby becoming familiar with the (your own) vocal tract. To that end, for example, students would probably not be tested on actual production/reception of speech in courses for which this book is used. Instead you will only benefit from knowing the definitions of kinds of speech events as outlined above. To this end a very useful for you to do is exercise 2-7 which asks you definitional questions about articulatory terms.

Table 2-2 Common English Words with Selected Sounds Described in Articulatory Terms*

" <u>p</u> ig"	aspirated voiceless bilabial stop	" <u>Ch</u> uck"	aspirated voiceless alveo-palatal affricate
" <u>c</u> at"	aspirated voiceless velar stop	" <u>z</u> oo"	voiced alveolar fricative
"pi <u>g</u> "	unaspirated voiced velar stop	"hi <u>ll</u> "	voiced alveolar lateral resonant
"lau <u>gh</u> "	voiceless dental fricative	" <u>w</u> indow"	voiced labialized velar resonant ⁷
" <u>h</u> igh"	voiceless glottal fricative	"s <u>i</u> ng"	unrounded high front vowel, nonnasal
" <u>n</u> ice"	voiced alveolar nasal resonant	"ma <u>sh</u> "	unrounded low front vowel, nasal

⁷ This is a co-articulated sound and could just as easily be called a velarized, bilabial resonant. In fact, in many charts the sound is presented in both the bilabial and velar columns.

“ <u>si</u> ng”	voiced velar nasal resonant	“Chu <u>ck</u> ”	unrounded mid central vowel, nonnasal
“ <u>y</u> es”	voiced alveo- palatal median resonant	“ni <u>ce</u> ”	forward glide from unrounded low central position, nonnasal
“wind <u>ow</u> ”	back glide from rounded mid back position, nonnasal		

*A more complete list of the articulatory configurations common in the author’s dialect of American English is presented in Appendix A.

Written Phonetic Description of Speech

The last part of this chapter deals with the way that linguists use this articulatory system to record speech. In other words, now that you have some idea of the positions of articulation and the manners in which these are used, how do you go about trying to describe what configurations are actually used by some speaker for a particular act of speaking. Some of the articulations can be seen and felt. Bilabial articulations for example, can be easily seen as can the lowering of the mandible for a low vowel. Strong aspiration can be felt by placing the hand just in front of the mouth as can voicing, assuming, of course, that for this particular articulation speakers will allow you to place your fingertips gently but firmly on their throat to feel the vibrations in the larynx.

But these visual and tactile clues, however useful, cannot substitute for the basic method of hearing the effect of different articulations. This obviously requires training so that you can hear a velar as opposed to a uvular articulation, or a mid back as opposed to a high back articulation. This isn’t as difficult as you may think, although it does require some time, since, as a human you are designed to hear the differences in hundreds of different articulations. You simply need to learn to apply the terminology of articulation to sound differences you (mostly) already hear. A “phonetician” is a linguist who specializes in describing speech sounds, and is typically adept at hearing a very large number of sound differences and accurately labeling these as to their place and manner of articulation.⁸

⁸ Again, I want to emphasize that no special physiological skill is required to be a phonetician, only training and practice in articulation identification. However, listening to, and recording human speech isn’t the only thing phoneticians do. They also study the acoustic properties of speech. They study the neurological and muscular activities of articulation (i.e., to investigate physiological parameters/constants of human articulation). Moreover, they can make use of other means of

Writing Speech Sounds

Let us now turn to the question of labeling speech sounds. We could do as we have done up to now and just use the position and manner terms in descriptive phrases, such as “voiceless dental fricative,” or “nasalized mid-front vowel.” This would provide explicit and relatively complete labels. Unfortunately, it is much too cumbersome and time consuming to be useful in describing such a rapidly occurring, and short-lived, phenomena as speaking. So over the course of the past 100 years or so, linguists have developed many specialized short-hand notations for speech (not for systems of writing which, of course, have been around much longer). The system presented here is a simplified system, and different in many respects from that presented by the International Phonetics Association (IPA). For a full IPA chart go to Appendix C. The reader is warned that there are different notation conventions, many of which have been in use for some time (American anthropological linguists used a somewhat separate notation for many years, some symbols from that system are also used here).

These systems traditionally draw upon the letters of the (Roman and Greek) alphabet for their graphs (each written symbol, representing one sound type, is called a phonetic **graph**). And therein lies a major obstacle in learning it: most of the graphs come already familiar to readers and therefore already burdened with one or more sound values. Consequently, students not only have to learn the articulatory meanings of each graph, but also unlearn their prior understanding of these graphs. Students in linguistics classes are warned that their ability to spell (which uses the prior references of the graphs) will be severely tested as they learn the new references. For example, the expected response to a joke is spelled “laughter” but when this word is spoken an articulatory notation would represent it as [l æ f t j]. Some of the graphs seem appropriate, such as [l] and [t]; some are unfamiliar, such as [æ], and one, [f], seems very wrong. This latter graph, and its use in “laughter” (that is, if you actually spelled “laughter” as “laufter”) would be considered a mark of ignorance and could be ridiculed, especially by an English teacher, if not by one’s more literate peers (although ‘texting’ conventions of condensed and alternate spelling may provide more room for such a spelling).

studying articulation such as placing a pressure sensing device on a speaker’s palate to record the placement and force of the tongue’s articulation, affixing an X-ray opaque material to the surface of the tongue and taking X-ray pictures of its shape, or even using optical fibers to light and photograph the activity of parts of the vocal tract such as the vocal chords in the larynx.

Many of the graphs will be unfamiliar and these will simply have to be memorized, such as [ð] or [θ] which represent the first sounds in “then” and “think” respectively. Nonetheless, the use of any articulatory, or “phonetic”, notation is filled with these kinds of dilemmas. To make matters worse, some speech segments are represented by phonetic notation exactly as they would be spelled in traditional American spelling. (The words “sip” and “met” would be [sip] and [met] so phonetic notation may not always look different.)

One System of Phonetic Notation

All of the current phonetic notations, including the one presented here, are historically created amalgams of many writing systems—as much a product of national pride (the predominance of Roman letters should come as no surprise given the Western European origin of most linguists) and convenience (graphs were often selected because they could be made on a typewriter keyboard) as by phonetic usefulness. Furthermore the notational graphs are extensively supplemented by smaller added marks called **diacritic marks**. These are dots, dashes, squiggles, chevrons, and even tiny graphs themselves placed to the left or right above, below and even through the main graph. (Most of the co-articulations briefly mentioned above in the expanded description of stops, pages 26—29, are indicated by diacritic marks.) The reader is warned that there is considerable variation in these supplemental marks. We do not need to concern ourselves with the complexities of a full phonetic notation. Instead the system presented here will show only some selected graphs, primarily those which have been useful in describing the articulatory configurations of English. A few others are added so that at least one graph is presented for each manner of articulation. Figures 2—9 and 2—10 have both the expanded consonant and vowel charts, described earlier in this chapter, but now with the addition of graphs to represent the speech sounds created by the interactions of place and manner of articulation. (For a more complete description of phonetic notation systems the student should refer to Abercrombie, Ashby, and Pullum and Ladusaw). It should also be noted that the system presented here may not be the same as the pronunciation guides given in many dictionaries. In the U. S., Anthropologists in particular, and linguists in general, have tended to use their own modifications of the International Phonetic Alphabet (IPA).

		PLACE OF ARTICULATION																	
		BILABIAL		LABIODENTAL		DENTAL		ALVEOLAR		ALVEOPALATAL		PALATAL		VELAR		UVULAR		GLOTTAL	
		vl	vd	vl	vd	vl	vd	vl	vd	vl	vd	vl	vd	vl	vd	vl	vd	vl	vd
		simple	p	b				t	d			k	g	k	g	q	g	ʔ	
MANNER OR ARTICULATION	STOPS	aspirated	p ^h	b ^h															
	affricates							ts	dz	ʧ	ʤ								
	FRICATIVES			f	v	θ	ð	s	z	ʃ	ʒ							h	
	RESONANTS	nasal	m					n	ɳ					ŋ					
	lateral							l											
	median	w								y		j							

Figure 2–9 Consonant Chart with Notation Graphs

Note: Shaded boxes indicate no sound of this type is possible. Partial list only.

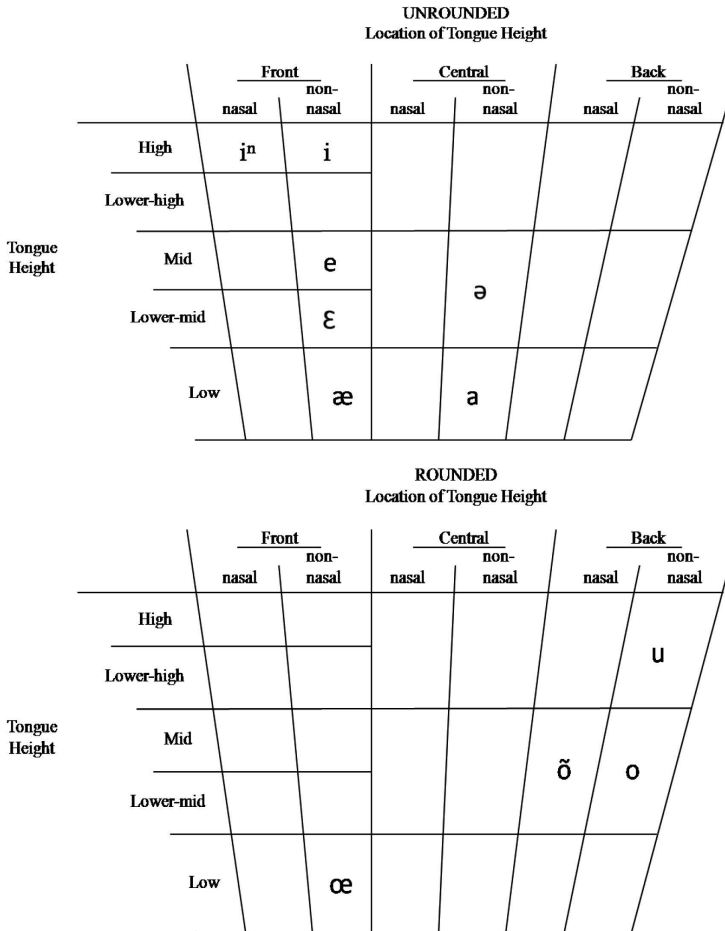


Figure 2–10 Vowel Charts with Notation Gaps
Partial list only.

Diacritic Marks for Graphs in Figure 2–9 and 2–10

For any consonant C: Ç = fronted from place for plain graph; Ç̣ = backed; C^h = aspirated consonant; C^ʔ = glottalized. For any vowel V: Vⁿ or Ṽ = nasalized vowel; V: or VV = long vowel; Vⁱ = unrounded forward glide; V^y = rounded forward glide; V^u = unrounded back glide; V^w = rounded back glide; V^ə = unrounded central glide. On Figure 2–7 glide 1 would be [oⁱ], glide 2 [aⁱ], and glide 3 [eⁱ]; Figure 2–8 glide 1 would be [i^ə], glide 2 [o^w], and glide 3 [a^w].

These graphs can now provide a shorthand for recording speech. For example, instead of writing: “voiceless aspirated bilabial stop; high front unrounded non-nasalized vowel; and voiced (unaspirated or simple) velar stop;” a linguist can describe the spoken word “pig” by just jotting down [p^hig]. (The brackets indicate that speech is being represented by phonetic graphs, quotation marks represent how a spoken form would be spelled.) Going back to the examples in Table 2–2 on page 34, here is how some of them as spoken by the author would be represented using this phonetic notation:

“cat”	[k ^h æ t]	“high”	[h a ⁱ]
“Chuck”	[č ə k]	“nice”	[n a ⁱ s]
“laugh”	[l æ f]	“sing”	[s i ⁿ ŋ]
“zoo”	[zu ^w]		

The advantages of phonetic notation as a way of recording speech are clear.⁹ But while it is necessary for students to have some mastery of phonetic notation to understand the kinds of linguistic analysis presented in the rest of this book, it must be remembered that notation is only a convenient means to a descriptive end, not linguistics itself. In other words, one doesn’t need to be a phonetician to appreciate linguistics. (But remember that accurate and reliable descriptions of actual speech events is indispensable to linguistic analysis.)

Each reader should try to pronounce each type of articulatory configuration and practice recognizing them in the speech of others (e.g. a friend). All that is required for the rest of the book is to know that [k], for example, is an un-aspirated voiceless velar stop or that [ð] is a voiced dental fricative, not to have to be able to produce these on demand or to correctly identify them in an informant’s speech. The actual field recording of an informant’s speech presents problems beyond the scope of this book and the student interested in learning more about this aspect of linguistics should start by consulting *Field Linguistics* by William J. Samarin (for an example of field recording see Appendix B).

Now try to apply your knowledge of phonetic graphs with exercise 2–7 (part A for the consonants and part B for vowels—if you wish you can just do both for each word). Give yourself plenty of room on the sheet(s) of

⁹ However, even skilled phoneticians need one or two repetitions of a single spoken word to record it precisely. Normal speech interaction, consisting of multi-word phrases, require the assistance of a tape recording played back many times. Even so, phonetic notation allows the linguist to quickly record some speech sounds (whether part of a long phrase or a single word) the first time it is spoken and then to fill in the blanks on repeated hearings. Refer to the example in Appendix B of a field record of speech forms.

paper you are using so you can go back if need be and make changes or add another graph. Toward this end, it is a good idea to use pencil, however many field linguists say that it is useful not to erase your “mistakes” since your mis-hearing of a sound can be useful.

Prosody: Stress, Tone and Syllables

There are two final, and very important, articulatory features to cover before we conclude this overview of phonetics. We have been discussing distinct articulatory configurations occurring within particular, very brief, time limits. A segment of speech would, then, be made up of a sequence of these configurations—different sounds occurring one after another associated with the corresponding succession of different configurations. But there are articulatory features which overlap, or extend over a sequence of individual configurations. If the single configurations are called “segmental” sounds, then what still needs to be discussed is what are called “supra-segmental,” or **prosodic**, sound features. These are stress and tone (sometimes called pitch).

Stress is the relative force exerted by the muscles associated with the vocal tract, especially the muscles around the lungs, usually manifested in speech as relative loudness. A higher stressed group of segmentary sounds will be relatively louder than adjacent groups of sounds with lesser stress. Stress is not bound to particular articulations—any configuration may occur with high or low stress, although particular languages may link certain sound types with certain stress levels (in American English, for example, the mid central vowel [ə] often occurs with low stress.) However, those sounds with less obstruction—resonants and vowels—and those which are voiced express stress more effectively. For this reason a resonant, and more particularly a vowel, often serves as being the center (or peak) of the degree of stress for a group of sounds. This stress peak is related to another important sound unit, the **syllable**.

A common definition of the syllable is a group of sounds surrounding, and including, a stress peak. Each syllable, consequently, will contain at least one resonant or vowel. Stress, by itself, is not sufficient to define a syllable (which is a challenge to define precisely), but for our purposes here—or, if you will, as a working hypothesis—we will assume an identity between stress and syllable. A single group of sounds, preceded and followed by silence, and containing one resonant or vowel is clearly one syllable. For example, the English word “bat” spoken in isolation is a syllable, and since there is no other adjacent group of sounds, we wouldn’t even need to assess its degree of stress. But the word “battle,” when spoken normally, contains two syllables since there are two different levels of stress at the two peaks—more on the first, less on the second.

“Battalion” apparently has three syllables since the stress level at the first and last peak is different (less) than the middle peak. In cases like this, one may not be too concerned with the relative difference between the stress levels of the first and last peaks since all that is necessary is that they are different from the peak which separates them. However, in a word like “corporation” the first two peaks have different stress but neither has as high a level as the third peak. Consequently a third stress level is required in order to use our definition of syllable and identify the syllable of the most stress. The three levels are called “primary” (relatively highest), “tertiary” (relatively lowest—usually referred to as no stress), and “secondary” (somewhere between primary and tertiary). “Corporation”, then, would exhibit a sequence of secondary, tertiary and primary stress levels for the first three peaks. It may be possible for a long utterance (i.e. a group of sounds which make up the words in a phrase) to have undifferentiated stress—i.e. the many vowels, nasals or approximants which would serve as peaks, all being produced with the same stress.¹⁰

In normal speech, however, stress levels would be arranged in different patterns, as much a part of the communicative signal as the arrangement of different segmental sounds. To fully describe speech, consequently, it is necessary to describe stress levels accurately. The delineation of syllables is one reason, especially for the linguist who wishes to describe how particular segmental sounds are combined into syllables. A more important reason is that stress can be used in a number of ways by speakers from word formation and the meaning of words to the construction of phrases and sentences and to signal the communicational intent of the speaker (asking a question, making a statement, or being sarcastic). Now try Exercise 2–9.

The other major prosodic feature is **tone**. We have already discussed voicing—voiced sounds include the vibration of the vocal chords in the larynx. The distinction between voiceless and voiced sounds, however, requires no reference to the relative frequency of the vocal chords’ vibration, either vibration occurs (= voiced) or it doesn’t (= voiceless). Tone, only applicable for voiced sounds, refers to degree of voicing. High tone refers to a group of sounds produced with voicing of a higher frequency than that of adjacent groups, which would have lower tone. This doesn’t refer to absolute frequency, or precise frequency. For example 200 vibrations per second, by itself, has no tone value. It may be interpreted as low if adjacent frequencies are higher, or as high if they are lower. Tone is

¹⁰ American English uses this arrangement to stereotypically depict the speech of someone under the influence of hypnosis or other ‘mind control.’ It can also appear in the speech of someone repeating a command with some exasperation - but in this case even though all syllables have the same (high) stress they would probably still be differentiated by pauses between each word of the command.

expressed only through voiced segmental sounds, unlike stress, but like stress, it is primarily perceived at the syllable peak, at a voiced resonant or vowel. Stress and tone are not necessarily congruent—a syllable with low stress may have high tone and vice versa. (American English does, however, combine the two so that high tone occurs only with high stress, and low tone with low stress. This makes it difficult for English speakers to hear these as different.) Also unlike stress, tone may be varied during a syllable peak. Some languages, for example Chinese, can distinguish syllables not by a difference in a steady rate of vibration during the syllable peak, but by the kind of frequency change: a Chinese syllable with the same segmental sounds as another but with a rising tone through the syllable peak will be different from the other having a falling tone, a dipping tone (falling then rising) or a steady tone. (Since American English speakers do not use tone in this way—where tone and tone “contours” over many syllables can be quite variable from sentence to sentence—it is also understandable why they would perceive a “sing-song” quality in the speech of a “tonal” language such as Chinese.) Like stress, tone must be described as an important part of the speech signal, and it can play the same range of communicational functions.

The notation for prosodic features presents a problem. First, it occurs over more than one segmental sound, so where is it to be marked? The method English dictionaries use for stress is to place a mark like ['] before a (written) syllable with primary stress. Of course since this assumes that the linguist already knows where syllable boundaries lie, it is not too useful in describing the speech of a language whose syllable boundaries have yet to be defined. One way of avoiding this problem is to mark the peak's stress. It is usual for American linguists to use the following notational diacritics (using [e] as an example): [é] = primary (highest) stress, [è] = secondary stress, and [e] (no diacritic) = tertiary stress (no or weak stress). If stress is marked over syllable peaks, then how will tone be indicated? A common method among American linguists is to use raised numbers as diacritic marks for tone. These can be placed anywhere tone seems to change. They are placed in front of segmental graphs (which carry the first tone or a new tone) and after the last sound of an utterance to indicate the final tone, if different. A four tone system seems to be useful as a start ([V] = any voiced sound), with [⁴V] representing the highest tone down to [¹V] representing the lowest tone. For languages, like Chinese, that have tone contours over syllables (i.e. rising, falling and rising), numbers such as these can be used with each representing one kind of contour. Try exercise 2–10 to test your ability to record and read phonetic notation.

With what I hope is the reader's dawning realization of the difficult task facing a field linguist trying to transcribe a sequence of segmental

sounds and sound features (e.g. aspiration, nasalization, voicing) as well as prosodic features, it is necessary to emphasize that a complete description is not achieved in one transcription of a single utterance. An informant/speaker is asked to repeat utterances several times with only some of the sounds and prosodic features noted each time. More likely a recording can be replayed many times and even slowed and/or stopped at particularly puzzling spots. (Once making a recording was at best a cumbersome process with all the attendant problems of a magnetic tape recorders and water, dust, insects, physical impact.) Some of these problems still remain but recording equipment as gotten much more compact and reliable (as was mentioned at the beginning of the chapter, laptops now can have software that not only record but do acoustic analyses. The field linguist as phonetician must be well-trained, but not super-human.

Terms to Know for Chapter Two

Phonetics	Vowel tongue positions
Acoustic phonetics	High Mid Low
Spectrogram	Front Central Back
Laboratory phonetics	Voicing
Articulatory phonetics	Aspiration
Manner of articulation	Affrication
Airstream	Nasal resonant
Ingressive	Lateral resonant
Egressive	Median resonant
Stop	Consonant co-articulations
Fricative	labialized
Resonant	glottalized
Consonant	Vowel co-articulations
Vowel	Rounding
Vocal tract	Nasalization
Tongue	Length
Tip	Glide (diphthong)
Front	Phonetic notation
Back	Phonetic graph
Root	Diacritic marks
Vocal cords (Glottis)	Prosody
Oral cavity	Stress
Nasal cavity	Tone (pitch)
Bilabial	Syllable
Labiodentals	
Dental	

Alveolar
Alveo-Palatal
Palatal
Velar
Uvular
Glottal

CHAPTER THREE

PHONOLOGY AND PHONEMIC ANALYSIS: SYSTEMS OF SPEECH SOUNDS

The preceding chapter dealt with a method of describing the very large number of human speech sounds. Although humans have the capability of producing thousands of distinct sounds, no speaker uses all of the speech sounds she is capable of making while speaking one language. In this chapter we will study that part of a language system which organizes (and limits) the speech sounds that individual speakers use. We will learn how to identify the set of distinctive kinds of speech sounds associated with each language. We will also identify how sound types—called phonemes—vary phonetically under certain conditions, and also learn how phonemes are organized into syllables.

The Phoneme Concept and the Requirements of Speech Communication

It should be clear that humans come equipped with the ability to produce a very large number of sounds. The purpose of the preceding chapter was to give the reader an introduction in describing the full range of human speech sounds based on how they are articulated. This should rightly be perceived as a challenging task, given the large number and variety of sounds produced and the complexity of the articulation process. However, the actual system of speech sounds employed by the speaker of one language is much less extensive than the range of human vocal capability. In fact, learning to be the speaker of a language is a process of learning how to categorize the large number of possible sounds into a relatively small number of distinctive sound types, called **phonemes**. The number of segmental phonemes varies from about 11 to about 70¹.

70% of all languages have between 23 and 44 segmental plus suprasegmental phonemes (Maddieson). Taking an average of about 30 phonemes, this seems a more reasonable number of sounds to worry about compared with thousands of possible sounds. While phonetics, then, is the study of all possible speech sounds by all human speakers, **phonology** (or

¹ !Xu, in southern Africa, has been reported as having as many as 141 segmental phonemes, including 48 “clicks” (see Maddieson).

phonemics) is the study of just the system of phonemes necessary to speak one language. That part of a language system which defines and organizes phonemes is called “the phonology of X language.”

Phonemes are sound types, or categories. It isn’t accurate to believe that if languages have on average about 30 phonemes, they would each have an average of only 30 sounds. Each phoneme category includes many sounds and so the total number of sounds regularly used by a speaker will run into the hundreds. So what is the difference between individual sounds (created by particular articulatory configurations) and phonemes? To answer this we must take a moment and discuss the process of speech communication.

Our speech sounds are signals used primarily to carry messages. Obviously, signals would have to be easily observable (hearable) and be able to be fairly readily distinguished from one another (otherwise groups of speech signals carrying one message might be mistaken for a similar group carrying another message). Also, we need a sufficient number of signals so that a sufficiently large number of different messages (constructed of different combinations of signals) can be used for the complex needs of human interaction. Let’s examine some hypothetical, but not very useful, types of speech sound signal systems in order to emphasize some salient characteristics of human phonologies.

Speech Based Upon a Phonology of Only Four Phonemes

First, suppose our phonology had only four phonemes: Phoneme 1 would be any kind of stop (bilabial to glottal, voiced or voiceless, affricated or not); 2 any kind of fricative; 3 any kind of resonant; and 4 any kind of vowel. These certainly pass the criterion of being readily distinguished from one another (with some minor uncertainties about affricates being confused with stop plus fricative, or nasalized vowel being confused with a nasal resonant). So, [pafk] (stop, vowel, fricative, stop) would clearly be a different total signal from [nafk] (nasal resonant, vowel, fricative, stop), and also different from [stip] (fricative, stop, vowel, stop). On the other hand, [kizb] would not be a different total signal from [pafk] (both are composed of a stop, vowel, fricative, and stop). So while different total signals can be made—and attached to different messages or meanings—there are limits to how many different groups of signals could be produced. Assuming that these four phonemes can be put together in any sequence, except that the same phonemes cannot occur side by side, we can calculate how many different total signals are possible: 4 signal groups consisting of one phoneme; plus 12 (4

$\times 3$)² of signal groups consisting two phonemes; plus 36 ($4 \times 3 \times 3$) with three phonemes; plus 108 ($4 \times 3 \times 3 \times 3$) with four phonemes; plus 324 with five phonemes; and plus 972 with six phonemes. The number of total signals consisting of one to six phonemes is no more than 1,456. (Of course, if we extended the size to ten phoneme groups, and even larger, the sum would be over 100,000, but experience with human languages has shown that basic total signals for each sound group with meaning—we are not talking about combinations such as compound words, or phrases—usually do not exceed six phonemes. For example, the average length of the basic signal groups in the preceding sentence when spoken is 3.5 phonemes.) And the total is actually much less than this since many of these phoneme groups would be impossible or, at least, unlikely.

For example: a nasal + stop + fricative + stop sequence (represented perhaps by [mpsk]) would be unlikely as would a stop + fricative + stop + fricative + stop sequence (such as [khkhk]). Nonetheless, even using the largest possible number of 1,458, this isn't even close to the number of signal groups required for human interaction. Think of each of these signal groups as having a message value, or meaning, and then try to imagine a human community getting by with only about 1500 basic meanings for all their environmental objects (animals, plants, landscape objects), manufactured items (tools, ornaments, ritual paraphernalia), natural categories (body parts, directions, cosmological entities, time, life cycle aspects), social groupings and relationships (marriage, family, relatives, friends, associations), and supernatural entities and forces, to say nothing of activities, modes of behavior and thought, objectives and evaluations. The reader can easily expand this short list, but it should be clear that 1,500 basic messages seems skimpy. It might be argued that 1,500 would be sufficient since these could contain multiple meanings and by combining them in clusters of two, three, or four or more a much larger number of meanings could be attained. It is certainly true that the essential and common meanings of a human community do not each have to be realized in individual basic signal groups. And, further, it might even be argued that there is some benefit to a language being "lean and mean"—i.e. being trimmed of all but the 1,500 truly essential, absolutely necessary meanings. The more meanings, one could say, the more confusion. Actually, there is no way of proving the error of either of those arguments or of proving the author's claim that any language needs about ten times the number of phonemes and at least ten times the number of signal groups. Except that linguistic research over the past two centuries has not

² The calculation is worked out as 4 times 3 since any of the four phonemes can begin but only three can come next since we will assume that the same phoneme would not be occur side by side.

come across any human group that speaks with this small a number of phonemes and stock of basic signal groups. In addition, consider that human communication is rarely perfect with the objects of discussion neatly defined—on the contrary it is more usually inconclusive, ambiguous and inefficient (if efficiency means saying all you need to, and only what you need to, with the smallest number of message units in the shortest time possible). Humans speak with background noise, interruptions, digressions, empty or polite phrases, and often enough without any precise “message” in mind but rather to sooth, fill a gap, cover embarrassment, or avoid silence. We do not seem to be capable of doing this with a minimal stock of signal groups.

There is also a technical problem with only using a small number of phonemes, each of which being a broadly defined category including lots of representative sounds, such as “all stops”. It places too great a demand on the users if the sounds can vary over great “articulatory” distances and thereby having a situation in which very different sounds can have the same effect. In the hypothetical situation described above, a [b] would function in the same way as a [ʔ] (both represent stops). Human speakers and hearers do categorize different sounds as having the same effect, but within much smaller limits so that all the sound members of a category would share some articulatory and acoustic features. For example, all bilabial stops as one phoneme in contrast to bilabial fricatives as another and alveolar and alveo-palatal stops as a third. The hearer doesn’t have to sort through a larger number of sound features to determine that a particular sound is, in fact, one member of a large diverse phoneme group, rather only a limited number of sound features, in some cases one or two, would be sufficient. Of course having the four phonemes narrowly defined would take care of this problem, but still leave the former difficulty of an inability to produce sufficient signals.

Speech Based Upon a Phonology of 150 Phonemes

The second hypothetical example postulates a phonology with 150 phonemes. Certainly this provides the basis for an ample number of basic signal groups (22,350— 150×149 —possible for groups of two phonemes alone). However, each of the 150 phonemes would have to be fairly narrowly defined since each would have to occupy much less articulatory ‘space’. For example, voiceless simple (unaspirated) stops could be grouped into eight phonemes by place of articulation (bilabial, dental, alveolar, alveo-palatal, palatal, velar, uvular, and glottal) and the same could be done with voiced simple stops for eight more (leaving out glottal). The same places could define sixteen more voiced and voiceless

aspirated stops, sixteen voiced and voiceless affricates, sixteen glottalized stops, and 6 voiced and voiceless laterally affricated stops for a total number of 71 stop phonemes. There would be a similar multiplicity of phoneme categories in the rest of the consonants and vowels. The problem would be that it would be a continual challenge for speakers and hearers to attend to the large number of very small differences between signal groups. Consider the following sequence of signal groups which differ from each other in only one or two articulatory features: [t̚ i k], [t̚ i k̚], [t̚ i k̚], [t̚ i k̚], [t̚ i k̚], [t̚ i k̚], [t̚ i k̚]; these could each carry different meanings since each is a different total signal. One might argue that this is often the case. Two signal groups that differ by a single articulatory feature and yet carrying different messages can't be that unusual. In English "kit" [k̚ i t̚] and "kid" [k̚ i d̚], for example, carry different meanings yet the only difference is the voicing of the final stop. While this is not uncommon in all languages, it can still commonly cause confusion: "Excuse me, did you say 'kit' or 'kid'?" Imagine the potential for confusion if utterances typically had several of these. One would have to speak very carefully, in conditions of no interfering noise, and hearers would have to pay attention to every sound—clearly not realistic. Human speech communication must succeed relatively well under poor conditions (speaker's less-than-perfect articulation due to injury—e.g. a lacerated tongue or a broken nose—fatigue or excitement as well as hearer's less-than-perfect reception due to noise, inattention or flawed speaker's performance) and this would not be possible for a system where lots of precise distinctions must be made in every utterance. An argument could be made that one wouldn't have to use words in the same utterance in which there is only one, very close, phoneme difference. To refer to our example above, with 150 phonemes it wouldn't even be necessary to use the rest of the groups similar to [tik] as having any meaning. If, say, [tik] meant "hot," then to avoid confusion, [tik] meaning "bottle" or [tik] meaning "drink," would not be used. Perhaps [gaz] meaning "container" could be used. There would be many other phoneme combinations to choose from, substantially different from [tik], to attach to meanings different from the one that [tik] has. As another example, if [tik] meant "hunk (as in an attractive sexy male)" then instead of having [tik^h] mean "geek," with the likely potential for misunderstanding, have [zam] mean "geek". The latter phoneme group is completely different from [tik] and it would be unlikely that the meanings "Jack is a hunk/geek" would be confused either by the speaker and/or the hearer. With 150 phonemes (having over 3 million possible three-phoneme combinations alone) this shouldn't be a problem.

However, the catch is that the purpose of a phoneme is to operate, by itself, in distinguishing one signal group from another. If having 150 phonemes results in avoiding certain phoneme contrasts, then there is no purpose in having the extra phonemes. Phoneme distinctions which are not used would undoubtedly disappear.

These two hypothetical examples hopefully illustrate that humans require enough phonemes to provide the means to construct a sufficiently large number of different signal groups but not so many that they get in each other's way or atrophy into extinction. This seems to be between about 23 to 44 phonemes.

Phonemes

A particular articulatory configuration, producing a particular sound, is not produced by a speaker, nor interpreted by a listener, in regard to all of its many acoustic characteristics, but rather for just those few (even only one) which distinguish it as a representative of one phoneme category and not any of several others close to it in "articulatory space." These sound characteristics are called **distinctive features**. Obviously we must speak by making particular sounds, but we organize and interpret these with reference to a system of phoneme categories. These phonemes are defined according to a few distinctive features, and it is the pattern of these features which linguists try to determine as they describe a language's phonology.³

The Phonemes of U. S. English

To get a better idea of a system of phonemes, let us examine some of the consonant phonemes of the author's dialect of U. S. English. There are only eight stop phonemes—six categories of simple stops and two categories of affricates. The distinctive features of these phoneme types are position, voicing, and manner (since two are affricates). Other articulatory features are present, of course, but these do not define different phonemes (these include aspiration and co-articulations such as labialization). In the bilabial area there are two phonemes, the distinctive difference between them (the phoneme boundary) is voicing. Regardless of whatever variations may occur for each phoneme, the versions of one are voiced and of the other voiceless. Common examples of these would

³ The articulatory/acoustic characteristics which are not necessary to define phonemes are not ignored, however. They are important in recognizing individual voices, making social judgments about speakers' backgrounds, interpreting speakers' intent (anger, humor, sorrow) and other assessments.

be the initial sounds of “bound” (voiced) and “pound” (voiceless), the middle sounds of “wiper” and “fiber,” or the final sounds of “cap” and “cab.” The reader should experiment with the range of variation possible with these phonemes. For example, try protruding your lips while making each one to see if the result suggests a different kind of sound (that is, not just an odd way of making the same sound, but a sound different enough to make the word different—in other words, does a labialization modification create a different phoneme.) Try varying the aspiration. Try making a labio-dental closure. This last variation should come too close to another phoneme boundary—not of another stop, but of the labio-dental fricative phonemes in English. Try a bilabial position but make a fricative (this will be a challenge since bilabial fricatives do not usually occur in English). This variation in manner should also produce a result which suggests different phonemes, again the labio-dental fricatives. In other words, there are articulatory boundaries to these bilabial phonemes. Moving to a labio-dental position is very close to, or actually over, one boundary of position and changing to a fricative manner is also very close to, or actually over another boundary. Within these boundaries are a number of articulatory variations which do not signal to a hearer that a different phoneme is being produced; these non-significant variant sounds are the actual sounds which constitute the phoneme categories of voiced and voiceless bilabial phonemes.

Within the area of “bilabial stop” voicing constitutes a boundary. Let us examine the nature of this boundary. If I stay within the bilabial stop area, but produce voicing upon the stop’s release, then I am in one of the two phoneme areas; if I do not produce voicing (i.e. upon release) I am in the other phoneme’s area. What is the significance of this change of manner. Simply put, it changes the signal sufficiently thereby alerting the hearer that there is a different message carried by the sound group. And this one change, by itself, is all that is necessary to change meaning. It is therefore proof of a phoneme boundary. For example, the two examples given above “cap” and “cab”⁴ ([k^hæp] and [k^hæb] respectively) clearly signal different messages. Assuming that the context was neutral (i.e. either of the following meanings was likely), so that if, for example, I said to my brother, “Richard, get me a Yellow Cab” it would be unlikely for him to go and get my yellow cap. I use “unlikely” deliberately since human communication is never guaranteed (Richard might be in a jocular mood). But any member of my community of speakers would be surprised if he did get a cap since we, Richard included, all recognize that a voiced

⁴ The voicing on final voiced stops in English is usually achieved by forcing a little air through vibrating vocal chords into the mouth even before the closure is opened. This is possible due to the somewhat elastic nature of the mouth walls.

bilabial stop is **not** the same signal as a voiceless bilabial stop. And even if the rest of the sound group, even the entire utterance, is identical, the voiced/voiceless change in the final position is sufficient to change the message. Contrast this with what would happen if I had said “Richard, get me a yellow [k̠^wæp], or [k̠^hæp^h], or [k̠^hæp^w] ([^w] represents labialization—protruding lips).” In each of these, although I might sound a bit odd, I, and any other member of my community of speakers would expect that Richard would return with a “cap.” Our community of speakers has learned that these variations (labialization, aspiration) are not sufficient to constitute a change in the signal (a voiceless bilabial stop phoneme) and therefore can not, by themselves, indicate that the message has changed. The next stop phoneme position boundary begins at the dental position and extends through the alveolar ridge to the palate. It is interrupted by the boundary of an affricated release, and these, in turn, as for the bilabials, are divided by a voicing boundary. Thus we have four more phonemes: voiced and voiceless simple dental/alveolar/alveo-palatal stops (which are usually just referred to as alveolar stops), and voiced and voiceless alveolar/alveo-palatal/palatal affricates (usually referred to as alveo-palatal affricates).⁵ Again, the reader is encouraged to try these sounds and experiment with varying them to where the sound produces a different signal. Examples of the voiced and voiceless simple stops, which also demonstrate that voicing produces a significant signal change, or a phonemic change, are “Tim” and “dim”, “Bertie” and “birdie”, and “hit” and “hid.” The change in manner (with some shift in position) to an affricated release also results in a phonemic change. Some examples of this difference, holding voicing constant, are “tuck” and “chuck” (voiceless), and “debt” and “jet” (voiced). The difference between the voiced and voiceless affricates can be illustrated with “jet” and “Chet,” “edger” and “etcher,” and “badge” and “batch.”

The last set of stop phonemes occur from the back of the palate through the velum to uvula (i.e. the very end of the roof of the mouth). These, as are all the other stops, are intersected by a voicing boundary so that there are a voiceless palatal/velar/uvular stop (usually referred to as a velar stop) and a voiced velar stop. Examples of these are, respectively,

⁵ The range for the simple stops is the back of the teeth (dental) to the front of the palate whereas the range for the affricate is further back—the alveolar to the back of the palate. This difference should suggest to you that phonologic systems (indeed all parts of a language) are not perfectly symmetrical. Some of the “imperfections” (as idealists would term them) are related to the constancy of change. Even assuming some tendency for consistency and balance some parts are modified in advance of or in different directions than others. However, the extreme complexity of language systems themselves would make it unlikely that even in the absence of change, complete uniformity would occur.

“core” and “gore,” “hackle” and “haggle,” and “pick” and “pig.” In contrast to the above stops, the reader should discover that a variation to palatal or uvular fricatives will not suggest different phonemes; that, for example, [xo^wj] ([x] is a voiceless velar fricative—not present in English) will be accepted as a version of “core” [ko^wj], or that [p^{hi} ɣ] ([ɣ] represents a voiced velar fricative) will be accepted as a version of “pig” [p^{hi}ɪg]. Unlike the bilabial and alveolar-palatal stops, which have fricatives articulated in the same region, and for which, therefore, a change to a fricative manner would cross a phoneme boundary, there are no palatal-uvular fricative phonemes in English and so the velar stop phoneme boundaries can extend across the manner feature.

The rest of the English phonemes are as follows:⁶

Fricatives—voiced and voiceless pairs of labio-dental (“file” and “vile”), inter-dental (“thigh” and “thy”), alveolar (“Sue” and “zoo”), and alveo-palatal (“rouge,” author’s pronunciation, some readers might say [r u^wʃ], and “rush”)⁷; and a voiceless glottal fricative (“high”); Resonants, nasal—voiced bilabial (“rum”), alveolar (“run”), and velar (“rung”); Resonant, lateral—voiced alveolar (“let”); Resonants, median—bilabial (“what”), alveo-palatal (“yacht”), and palatal (“rot”); Vowels, simple—high front unrounded (“bit”), mid front unrounded (“bet”), low front unrounded (“bat”), mid central unrounded (“but”), low central unrounded (“bought”); Vowels, glides—front glide from high front (“beat”), front glide from mid front (“bait”), front glide from low central (“bite”), front glide from mid back rounded (“boy”), back glide from low central (“bout”), back glide from mid back rounded (“boat”), and a back glide from a high back rounded starting position (“boot”). This gives a total of 36 segmental phonemes.

While prosodic, or suprasegmental, features play important roles in the use of speech to social and personal ends (joking, questioning, sarcasm), there seems to be only one prosodic feature in English which can change the basic signal and this is stress (typically accompanied by tone). “Permit” [p^hejmít] is a verb with high stress on the second syllable and “permit” [p^hejmít] is a noun with high stress on the first syllable.

⁶ For simplicity’s sake just the conventional (or, as some would term it, the target) articulation is given, not the full range within the phoneme category. For example, the term “alveolar” stop phoneme is used for a phoneme category which extends from dental to a front palatal position. Only one example is given for each, usually in an initial position.

⁷ Note that the vowels in “rouge” and “rush” are not the same. It is not an easy task to find English words which only differ in the voicing of this fricative.

General Characteristics of Phoneme Systems

Note that English phonemes represent the range of articulatory configurations and processes. Consonants occur from bilabial to velar and glottal positions, and from stops to resonants. Vowels occur from front to back, high to low, and as rounded or unrounded, simple and complex (glides). Certain areas or features are not significant (for example, lateral or palatal fricatives, glottalization, nasalized vowels) but even so, if one thinks of all possible articulatory configurations as a large space, then all general regions of it are represented in English although all of its many parts may not be. Moreover, this condition is true for all languages—all regions of articulatory “space” will be represented. In fact, we can now state the following assumptions concerning the phonologies of different languages: 1) all regions (if not all specific parts) of articulatory space will be represented, 2) the many speech sounds used by speakers will be organized into distinctive phoneme categories within which articulatory differences are non-significant, 3) there are a small number of distinctive articulatory features which separate phonemes, and these features are patterned (as will be demonstrated below), and 4) the system of phonetic notation will be sufficient to describe all sounds (phonemic as well as non-phonemic) used by human speakers. This last assumption provides the starting point for the linguistic analysis of a language’s phonology: a (fairly)⁸ complete and accurate description of valid speech activity.

Phonemic Analysis

Recording a Corpus

For the sake of simplicity, the initial phonetic record is done on very short speech items (i.e. such as single “words”). Obviously this eases the phonetic task considerably. A linguist can handle the description of “chair” [tʰej] (even including many repetitions and corrections/additions) as a response to “what is that?” but would have some difficulty with a response like “That there? Well that’s my father’s favorite chair, the one his grandfather made and hauled all the way over from the old country.” However people rarely interact in single word responses and so the strategy of working through a list (“What is this,” “now what is that,” “and what is this”.....) undoubtedly produces non-natural speech. Still, as long

⁸ An absolutely complete description, even if it were possible, is not necessary for the reason that the analytic method to be described will reveal any gaps in the initial phonetic record and so it is best to get on with it rather than taking time to get all possible phonetic details.

as the results are valid speech forms, which could potentially occur in isolation as an appropriate speech activity, the benefits provided by the ease of phonetic recording outweigh the disadvantages of getting, for the most part, non-typical speech. (There is another reason for trying to limit responses to brief speech units which will be mentioned below.)

The objective is to obtain a large number of speech segments (at least several hundred) recorded phonetically along with some assessment of their meaning. The whole list is called a **corpus** (“set of speech data” is a more accurate term but too long), each phonetic description is called an **item**, usually given some identification (a number, e.g. 16.52 for item 52 in list 16—body part terms), and each meaning (really an estimate translated into the linguist’s language) is called a **gloss**.⁹

Now, the reader should have anticipated that the linguist’s phonetic recording will include lots (hundreds) of articulatory features which play no role in distinguishing one phoneme from another, as well as the handful of features which do. *The question becomes, which is which?* For the language being studied, which sounds can be grouped together as variants of the same phoneme? Which are the distinctive features which distinguish its phonemes (separate them from one another)? While these may seem to be formidable questions, in fact, an outline of most of a language’s phonology can be worked out in a remarkably short time. Phonologies cannot be highly abstruse, intricately enigmatic puzzles for the simple reason that they must be readily useable by their speakers. This means that they are also readily interpreted by linguistic analysis.

Minimal Pair Analysis

The most straightforward means of establishing phoneme boundaries is by using **minimal pair** analysis. Many of the examples given above in the overview of English phonemes were minimal pairs. These are valid speech segments having clearly different meanings (i.e. they carry different messages), and are composed of identical sequences of sounds except that there are different sounds at the same location in each. The difference(s) between these two sounds, therefore, must be sufficient to signal that each sound group carries a difference in meaning and so must be a distinctive, or phonemic, sound feature.

⁹ For an excellent account of what a good, usable corpus should be like, see Samarin. In this book we will be using “modified” corpora. Not only are they very small (real corpora contain at least hundreds or, more likely, thousands of items), but, more importantly, they are very selective—having just those items which illustrate a point in the book, or pose a certain problem or analytic technique. Real corpora are not this accommodating. Refer to Appendix B for an example of a corpus and how it is obtained.

Actually a single minimal pair is not too useful. It only demonstrates that there is a phoneme boundary between the two different sounds. It doesn't indicate any of the other boundaries of the two phonemes nor does it indicate the non-distinctive variation within the phoneme (what is called **allophonic variation**—see below). Further, it may not even immediately reveal the actual phonemic boundary. For example, for the English alveolar stops, the supposed minimal pair “Tim” and “dim” ([t^him] and [dim] in the author's dialect) actually represent two possible boundaries, one of which is the voiceless-voiced boundary which, as mentioned above, turns out to be the actual phonemic boundary. However, there is also the difference of aspirated – unaspirated. But, and this must be emphasized, a linguist just beginning to analyze English phonology (one who didn't know English) would not know which of these differences was phonemic. Perhaps the aspirated—unaspirated difference could be the basis for a phonemic boundary with voicing being non-phonemic. By itself this minimal pair would not provide sufficient data to determine the actual boundary. Of course the linguist could also assume that there are two boundaries consisting of both features (in other words four phonemes—a voiceless aspirated, a voiceless unaspirated, a voiced aspirated, and a voiced unaspirated phoneme). There is nothing to do but hope to obtain more minimal pairs involving this type of sound—in particular one in which the alveolar stops occur in different locations (i.e. other than in an initial position). Suppose the linguist happens to obtain “hit” and “hid” (again, in the author's dialect, [hit] and [hid] respectively). After first making sure that these two items have different meanings, the linguist now has a minimal pair in which the only difference is in voicing. More minimal pairs could be hoped for, but this would seem to support the (tentative) identification of voicing as one of the phonemic boundaries among alveolar stops, and not aspiration.

She might go back to the first pair and try saying them to her informant with the aspiration and position identical and only the voicing different to see if the informant will accept them as different to test this tentative conclusion. However, using yourself (a non-native speaker) as an “informant” presents many risks, one being that the real speaker might simply assume that you made an error in pronunciation by leaving out the difference in aspiration (which could be important) and, rather than be impolite, agree that they are different.

Even if voicing can be isolated as a phonemic boundary, the linguist must still determine other boundaries. Let us suppose she obtains “hip,” “hiss,” and “hick” ([hip], [his], and [hik]). Putting each of these together with just “hit”, she will have three more minimal pairs which (assuming for the moment that the problem of multiple possible boundaries described above does not arise) will yield: 1) a boundary of manner (stop—

fricative), with the voiceless alveolar fricative ([hit]—[his]), since both are voiceless and unaspirated; 2) a boundary of position (alveolar—bilabial) with the voiceless bilabial stop ([hit]—[hip]), since both are stops and voiceless, and 3) another boundary of position (alveolar—velar) with the voiceless velar stop ([hit]—[hik]). As more minimal pairs are obtained for different sounds then more boundaries will be determined.

But these boundaries must not be taken as anything more than estimates until the entire phonemic system is determined. For example, suppose our linguist obtained “here” ([hiɹ]) as an item but no others which could be the other part of a minimal pair with [hit] or [hid] (i.e. she would not have obtained [hip], [his], or [hik]). She would indeed have what would appear to be two minimal pairs—[hit] and [hij], and [hid] and [hij]—but to what gain? It is of no real use to prove that there are phoneme boundaries between [j] (a voiced palatal median resonant) and [t] and [d] (voiceless and voiced, unaspirated alveolar stops). They are just too “distant” in articulatory space to be members of the same phoneme. Contrast the minimal pair, “hit” and “hid”, which proved that there was a boundary between the closely similar, unaspirated voiceless and voiced stops. This was a very useful finding because [t] and [d] might very likely have been members of the same phoneme. In fact, we can assume for now that sounds with more than two articulatory differences probably are not members of the same phoneme. This doesn’t mean that any linguist will refuse or discard items which create minimal pairs for widely different sounds, but it does suggest priorities—our linguist might try, if it were possible, to elicit items involving sounds different only in voicing or aspiration—sounds which would be different in just a single articulatory feature.

Exercise 3–1 provides you an opportunity to identify minimal pairs in English. In this exercise you will have to play around with the sounds of words (as you might do if you were trying to rhyme a word for a song or poem), but you can rest assured that for almost all of the sounds minimal pairs exist—at least in one of the positions (initial, medial, and/or final).

Actually, minimal pairs are not really high on any linguist’s research agenda. There are several reasons for this. One, of course, mentioned above, is that they may not provide much useful information. More importantly, they may occur very infrequently among the items in a corpus. English has a large number of minimal pairs—but English word structure allows for a large number of one and two syllable words. The smaller the word, the greater chance there will be of finding another which differs only in one sound. Many languages have words made up of many sub-parts (these will be called “morphemes” in chapter 4) and while the sub-parts themselves are small, they never occur by themselves such that they could serve in a minimal pair. In any event, the words formed by their

combinations are longer and it is difficult to obtain minimal pairs for words of about ten phonemes. Try to think of a minimal pair with two five syllable words (remember, there can only be the one difference, and no others.)

Distributional Analysis

There is another approach to determining phonemes which doesn't rely on minimal pairs. We will call it **distributional analysis**. It can lead to a complete description of phoneme boundaries as well as the variation within each phoneme. It is based on the premise that the sounds which make up a phoneme category possess two attributes: 1) they are phonetically similar; and 2) they have a distinctive pattern of occurrence in regard to their sound environments.

To state that they are phonetically similar is to say in effect that they are in the same region of 'articulatory space.' Since the items in a corpus are represented using phonetic graphs, you will have to have some idea of what kind of speech sound each phonetic graph represents (place, manner, voicing, nasalization....) or you will have to keep looking each graph up in the phonetic charts on pages 38 to 39, or in the more complete phonetic charts in Appendix C, in order to identify those graphs in the corpus which represent similar sounds. Having studied phonetic notation at the end of chapter 2, you should be able to look at two or more graphs and tell if they are different in only one, or at most two, phonetic details or not. Distributional analysis is only performed on sounds which are phonetically similar—sounds which might possibly be members of the same phoneme.

At this point you should try to apply what you have learned so far by doing exercises 3-2 and 3-3. Both present you with a corpus and you task is to record the sounds represented in the corpus (vowels for 3-1 and consonants for 3-2). When you are done with your chart, examine it and circle those phones that are close enough in articulatory space that they might be members of the same phoneme.

Once similar sounds are identified from the phonetic corpus (these are called **suspicious pairs**—see below), then each sound's environment is listed. By "environment," we mean those speech sounds which occur immediately preceding and immediately following the sound we are examining. This is the contiguous sound environment, those sounds (or articulatory configurations) just prior to the sound we are examining and just following it. According to our operating premise about the sounds which are members of the same phoneme group (we will now call these by their linguistic name—**allophones**), they will not occur randomly, nor will they occur in identical preceding and following phonetic environments. Instead, each allophone occurs in a particular range of phonetic

environments (in contrast to the environments of the other allophones of the phoneme). In other words, the total environment that the whole phoneme group occurs in is divided up among the member allophones so that their individual environments ‘complement’ each other—i.e. the sum of all the allophonic sub-environments equals the total environment of the phoneme. This is termed **complementary distribution**. This means that if the linguist can determine that 1) two (or more) **phones**¹⁰ are articulatorily similar, and 2) only occur in different environments then he may conclude (initially) that they are allophones of the same phoneme. Conversely, since phones which are allophones of different phonemes would typically occur in the same environments¹¹, then if two phones have the same pattern of occurrence (they have identical preceding and following environments) then the linguist may conclude that they belong to different phonemes.

Let’s look at an example, again from the author’s dialect of English. Suppose that the author served as an informant for a linguist studying English and thus provided a number of items for a phonetic corpus. It is likely that both a [k^h] and a [k^h] (remember that the first graph indicates a palatal stop and the second a velar) would occur in some of the items, and therefore, because there is only one phonetic difference between these two aspirated voiceless stops—place of articulation, one is palatal and the other is velar—the linguist would consider these two phones as possibly members of the same phoneme (because they are phonetically very similar they would be a suspicious pair). If a minimal pair was discovered in the corpus with these two phones being the only phonetic difference, then the matter would be settled in favor of them belonging to different phonemes. The existence of a minimal pair (for the sake of our illustration suppose these were [k^hos] and [k^hos])¹² would in fact mean that the two phones can occur in exactly the same preceding and following sound environments, in this case both would be able to occur at the beginning of an item—preceded by silence—and followed by a rounded non-nasal mid-back simple vowel. Consequently they would have to be members of different

¹⁰ During analysis, when the linguist is working with individual sounds and not yet sure if they are members of different phonemes or allophones of the same phoneme, sounds are just referred to as phones.

¹¹ This isn’t completely true for all the phonemes of a phonology, i.e. vowels probably won’t occur in all of the same environments that consonants will, although there will be some overlap. But a vowel and a consonant are not going to belong to the same phoneme. However, the phonemes which are of the same type (e.g., all stops, all glides), probably will share pretty much the same environments.

¹² Actually neither of these items are real—i.e. neither are in the author’s dialect. They are only hypothetical illustrations and typically would have been marked by an asterisk to indicate that they are “unreal.”

phonemes since they would not satisfy the second condition of being an allophone of the same phoneme—complementary distribution. But let us suppose that no minimal pair occurs (or, more appropriately, the linguist doesn't look for any, relying instead on distributional analysis.) The linguist would start listing the preceding and following sounds of each occurrence of [k̤] and [k^h] in the corpus. It is necessary to do this work carefully since missing an environment might cause the linguist to come to an incorrect conclusion. Let us further suppose that our corpus includes items such as [# k̤ i#] “key”, [# k^hət#] “cut”, [#k^hátə#] “cotton”, and [#k̤ itin#] “kitten” (the number symbol, #, indicates silence, which would be what would precede, and follow, an item said in isolation). Listing the environmental data for each phone would look something like the following:¹³

#		i ⁱ	#		ə t
#		k̤	#		k ^h
#		í t i ⁿ n	#		á t ə ⁿ n

Obviously, an actual analysis would depend upon more data than just two occurrences of each phone, but let us suppose that even if we had many more items containing either of them we would still have the same information. The linguist can now see that both phones can occur preceded by silence [#] and therefore there is an identity in the preceding environment. However there are **no** identities in the following environment [k̤] occurs followed by [iⁱ] and [i] while [k^h] occurs followed by [ə] and [á]. Therefore, because they do not occur in the same sound environments, they satisfy the second attribute of being allophones of the same phoneme—*occurring in complementary distribution*. In this particular case, it is the following sound environment which is completely different. (We should add here that when allophones occur in complementary distribution it is usually only in one environment, either the preceding environments are completely different **or** the following environments are. It is not necessary for both to be different, and, in fact, this would be unusual. It must also be noted that it is not necessary for all of the sounds in an environment to be identical in order to have an identity in the environment; to say that one of the environments is “identical” is to say that there is (at least) **one** identity in the environment. In other words the two phones *can occur in an identical environment*, but this is not

¹³ There are several ways of organizing the presentation of phonetic environments; this is simply the way the author prefers.

necessary in all cases of occurrence. (Perhaps a better way of stating this condition is to state that “the preceding/following environments **contain** one or more identities.”) In our example above we only had one sound type in the preceding environment—silence [#]—but do not be misled into thinking that this will always be the case. Nonetheless, since only one identity might be found, this underscores the need for careful work.

In our example, however, the following environments of [k^h] and [k̥] have **no** identities. Consider what this means. These two phones can never occur in a minimal pair. Potential minimal pairs such as [k̥i^ht] and *[k̥i^ht] or *[k̥ə^ht] and [k̥ə^ht] (remember that the asterisk indicates an impossible item, not a real item) could not occur since [k^h] could not occur followed by [i^h] (in *[k^h i^h t]) nor could [k̥^h] occur followed by [ə] (in *[k̥^h ət]). Not being able to occur in identical environments means that there is no way that changing from one of these phones to the other could serve as a signal change in an otherwise identical sound context and thereby be the basis for a change of message. In other words, the difference between [k̥^h] and [k^h] cannot be a distinctive sound change. ***The difference between them is not a phoneme boundary. Instead, the difference is allophonic.*** It represents a variation of the same sound. These two phones must be considered as variants of the same phoneme—two allophones of the same phoneme.

Now let's look at another example. Our linguist also notices that [k^h] and [g] (which also occurs in the corpus) are similar enough that they too might be members of the same phoneme. And again, not bothering to look for minimal pairs, carefully notes the environments in which [g] occurs ([k^h] has already been done). [g] occurs in [gəm] “gum,” [gəs] “Gus,” and [gad] “God,” and so the environmental data will be as follows (the data for [k^h] is presented again for ease of comparison):

#		k ^h		ə t		#		g		ə m
#				á t ə n		#				ə s
										a d

We can see that both occur preceded by silence [#] and both occur followed by [ə]—i.e. they can occur in identical environments (we must ignore the [á] phone since one is stressed and the other is not). Our linguist could look for a minimal pair since it seems that it is possible. We could even suppose that one turns up several weeks later when /k̥ə^hm/ (“come”) is obtained, with a different meaning than [g əⁿ m] (“gum”). The important point here is that the linguist really doesn't have to deliberately search for a minimal pair. The fact that both [k^h] and [g] can occur in the same

environment indicates that they can be the basis for a significant signal change, and that they therefore are certainly members of different phonemes.

Basic Strategy of Distributional Analysis

We can now state the basic strategy of distributional analysis. Starting with a fairly complete phonetic corpus (in other words, the linguist has done a workmanlike job of accurately representing each item in the corpus using phonetic notation), the linguist: 1) constructs phonetic charts and places all of the phones in the corpus on them; 2) determines which of these are phonetically similar¹⁴; 3) taking each group of similar sounds in turn, lists the preceding and following environments in which each of these phones occur; 4) examines the environments to see if there are identities in the preceding environments of the phones and then in the following environments; and 5) states conclusions about their *phonemic status*. Those pairs with one of the environments (preceding or following) *completely* different (therefore being in complementary distribution) are labeled as allophones of the same phoneme, and those pairs with at least one identity in both preceding and following environments (at least one exact match in their preceding environments and at least one exact match in their following environments—the whole environment does not need to consist of identical sounds) are considered to be members (allophones) of different phonemes.

Now you should try exercises 3–4 through to 3–8. Each of these is a problem in determining whether or not two phones (which are similar, but you should always check to make sure) are allophones of the same phoneme or not. You will use the method described above as distributional analysis to figure out if the two phones occur in complementary distribution.

It isn't as simple as the above description suggests. For one thing, the phonetic transcription may contain erroneous additions and omissions of phonetic details. The linguist, after all, even with all of his training, is still a creature of his own language's phonology and may listen "with his own ear" from time to time. An English-speaking linguist, for example, could believe he "hears" a [k^h] before a central vowel even though the informant actually says a [k^h] (because this kind of allophonic occurrence is ingrained into his English-speaking mind). The informant, as well, might

¹⁴ Phones which are phonetically similar, and which are therefore potential candidates for inclusion within the same phoneme are traditionally called "**suspicious pairs**" even though there may be more than two. They are "suspected" of being allophones of the same phoneme.

produce what, to a fellow speaker, would be an unusual articulation. This is likely to occur when speaking overly slowly for the linguist's benefit or toward the end of a tiresome session of repeating simple words for an odd stranger who can't seem to say them back correctly. However, these conditions just point up the fact that linguistic research proceeds along with the collection of speech samples. The linguist (like an ethnographer) doesn't just collect a mass of "(speech) facts" and then afterwards try to make sense of them. Instead, analysis goes along with data collection. Instances such as the incorrect hearing of [k^h] referred to above will show up as anomalies before central vowels and the linguist can recheck the items in which they appear.

Allophonic Conditioning and Allophonic Consistency

There are two important characteristics of phonologies which assist this rechecking process. The first concerns the system of allophonic variation within phonemes and the second concerns the systematic nature of phonologies. Allophonic variation is not random. We have seen that the allophones of a phoneme are in complementary distribution, but this shouldn't be taken to mean only that they merely don't get in each other's way. The allocation of allophones is related to kinds of articulatory modifications which speakers have learned to make on their phonemes. Moreover, these modifications are not unique to each phoneme but tend to be constant over whole sub-sets of phonemes (e.g., all stops, all front vowels). Let us look at the example of the English [k^h] and [k^h] allophones of the voiceless "velar" stop phoneme. We saw above that

[k̤^h] occurs followed by front vowels while [k^h] occurs followed by central vowels. We demonstrated this with only a few examples, but let us suppose that this distribution had been firmly established by many examples in our corpus. The linguist would note that this is unlikely to be an isolated allophone distribution—only occurring for this phoneme. Instead she would expect that it would be characteristic of some sub-set of phonemes: all velar phonemes, all voiceless stops, all velar stops, or even all stops. Upon further analysis, it will turn out that this distribution doesn't occur for the allophones of bilabial and alveolar stop phonemes, but it does occur for voiced velar stops and also for (voiced) velar nasal resonants. In other words, in English, velar phones are "conditioned" by following vowel position such that palatal and velar positions are allophonic variations (as well as uvular, but we are not considering these now). Now how could this help our linguist's analysis, in particular in regard to checking insufficient or even erroneous data? Suppose that our linguist had adequate data on velar nasal phones which showed this

distribution but that the data for the voiced velar stop was insufficient. For example, the only occurrences of voiced velar stops were of [g] and even though these were indeed followed by central vowels, this would be insufficient evidence of complementary distribution without [g]. On the basis of the assumption of allophonic consistency she would tentatively conclude that if voiceless velar stops and voiced velar nasals had this allophonic distribution, then voiced velar stops (and any other velar, for that matter) probably had it as well. This isn't proof, certainly, but it can help focus the analysis. Or, suppose that there was one instance of [g] in the corpus and it was followed by a central vowel. This would raise the linguist's suspicion that perhaps it was an error and should be rechecked. Again, the consistency principle of distribution is not proof, but consider how challenging phonemic analysis would be if each phoneme had its own allophonic principles, different from all other phonemes.

Assimilation

It was mentioned above that allophonic variation can be considered to be the result of conditioning by type of environment. Let us pause for a moment and examine the way this conditioning works. One form of conditioning is referred to as **assimilation**. This is where an allophonic variation exists because it is the part of the phoneme category most like the environment. For example, the English allophones [k̟] and [k̠] (see above) are the result of assimilation in the following manner: [k̟] is a fronted (palatal) version of the phoneme (which, remember ranges in place of articulation from palatal to uvular) and [k̠] is a central (velar) version. [k̟] occurs followed by front vowels and [k̠] by central vowels. Therefore one could argue that when the following articulation is forward in the mouth then the velar stop is assimilated to this "forwardness" and therefore made in the forward range of the phoneme. In fact, assimilation is often explained as being based in articulatory efficiency. In the case of [k̟], for example, since the tongue will be forward in the mouth for the following vowel then it is more efficient to make the closure for the stop in a forward position, and the [k̠] allophone can therefore be similarly explained as being made further back because the following central vowel has the highest part of the tongue further back in the mouth. I will leave it to the reader to think through the flaws in the above "efficiency-based" explanation, and, indeed, in any attempt to base allophonic variation just on physiological assimilation. Without doubt there exist combinations of articulatory configurations that lend themselves to a merging of position and or manner, but these must be learned. To say they are physiological

must imply that they would occur for all humans. Becoming an appropriately skilled speaker means that among all the other skills, one must learn how to make allophonic variation in the appropriate environments.¹⁵ And different speakers using different phonologies learn to respond allophonically to environments in different ways.

Now you should attempt exercise 3–9. We will return to Zoque (used in exercise 3–2) and see if we can determine a consistent pattern in Zoque allophonic variation for stops.

Phonemic Consistency

The second principle which can assist the linguist's analysis of phonology is **phonemic consistency**. Similar to the principle of allophonic consistency over sets of phonemes, phonemic consistency refers to use of the same phoneme boundaries among sets of phonemes. In English, for example, voicing is a phoneme boundary for each of the stop and fricative manners (with the exception of the glottal fricative, which can only be voiceless). Consequently, the linguist who had good distributional evidence for English voiced and voiceless bilabial and alveolar stops but inconclusive or even (erroneously) contradictory evidence on voiced and voiceless velar stops (i.e. that they appeared to be in complementary distribution) might look carefully again at these velars, or at least not come to a conclusion until more evidence is obtained on the assumption that voicing would also be phonemic for velar stops if it is for the other stops. The principle of phonemic consistency is, again, not proof, and would certainly not take the place of valid phonetic data in the corpus, but it can focus the analysis.

Phonemic Description

Phonemes and Their Allophones

The result of phonemic analysis will be a list (or chart) of the phonemes in the language's phonology. Each phoneme's allophones will also be described along with their distribution (or, if presented as a chart, the boundary allophones for each phoneme). Since the allophonic environments are easier to describe in a list format, this is the way we will

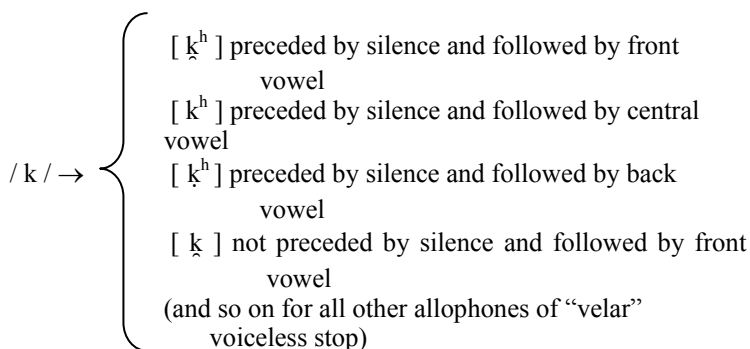
¹⁵ Actually one clue that a person is a non-native (but understandable) speaker is if he or she does not make the appropriate allophonic variations. The phoneme boundaries must be respected, of course, or the person will not be understood, but failure to appropriately modify phonemes is part of the 'distortion' that gives the non-native speaker an "accent."

use here. The list would include entries something like this: (from English, using the voiceless velar stop phoneme):

General Phoneme		Environment
Description	Allophones	for Allophones
voiceless palatal—	[ɰ ^h]	preceded by silence and
velar—uvular stop		followed by front vowel
	[k ^h]	preceded by silence and
		followed by central vowel
	[k̠ ^h]	preceded by silence and
		followed by back vowel
(other allophones not included)		

Phoneme Graphs

Each of the other phonemes would also be listed along with their allophones and allophonic environments. However, rather than describe the phoneme using its general articulatory features, the linguistic convention is to use a phonetic notational graph. In the example given above, instead of describing the phonemes as a “voiceless palatal—velar stop”, just the graph “k” would be used. Quite frankly, there is a problem in doing this. Phonetic graphs are best used to record articulatory detail, not general phonemic categories. In fact the graph “k” doesn’t properly indicate a “palatal-velar stop”, but really only a “simple velar stop, weakly aspirated”. To offset this problem somewhat, graphs used to represent phonemes are enclosed in slant lines or slashes (/ /), Thus the above example would be represented as follows:



The graph in the slashes stands for the phoneme, the arrows represent “is allophonically manifested as,” and the graph in brackets stands for the particular allophone which occurs in the environmental conditions listed afterwards. The advantages of using phonetic graphs in this double fashion—to represent actual sounds (allophones), as well as to represent the phoneme category—is somewhat confusing, but for now we will do this to streamline the listing of phonemes (since they are categories of speech sounds).

Speech will be analyzed for grammatical and semantic patterns in addition to phonological patterns and these analyses are done on phoneme sequences not allophone sequences. Remember that allophonic variation does not make any difference in meaning and therefore has no role in the analysis of grammatical/semantic units. A form like [k^hik] (“kick”) in English would not need to be represented using the two allophones of /k/ when, for example, it is being used grammatically in “kicked”—the past tense inflected form. So instead of [k^hikt] (written using phonetic graphs) we would write /kikt/ (using phoneme symbols). Phonetic details would get in the way of grammatical and semantic analysis as overly accurate “clutter.” In fact the speakers themselves, whose language the linguist is trying to describe, don’t usually need to be aware of allophonic variation. As far as an English speaker is concerned both allophones of /k/ in [k^hikt] “sound” the same. (It is thus always with some difficulty that the linguistic teacher convinces their students that the allophones they use as variants of their phonemes are, in fact, different sounds.)

Deciding on Phoneme Notation

Selecting which phonetic graph will serve to represent the phoneme presents some problems. Technically, the graph for the allophone which has the most frequent occurrence or the one which indicates the common features of most of the allophones should be picked. (In our English velar phoneme example, none has a greater distribution, but */k^h/ does indicate aspiration, which three of the allophones possess.) However, convenience of use (in writing and, especially, typing) is taken into consideration. And since phonemic representation must not be so technically accurate that it stands in the way of clearly and efficiently written forms, diacritic marks are often discarded, if possible. Thus the use of just the /k/ graph.

There is another consideration, one that is political. Since phoneme symbols do not represent actual sounds, only categories of sounds, one could use any mark, even numbers, for the phoneme categories. So regardless of a phonemic analysis that is done using the kinds of phonetic graphs shown here, what might be selected by a society for its phoneme

symbols would be the written symbols of the dominant national group of which they are a member. The velar phoneme discussed above in English, could be represented by /x/ “chi” if English were spoken in a context where the dominant written system was the Cyrillic alphabet (e. g. Russia). Or, it might be represented by /خ/ “kha” if the dominant system was Arabic (e. g. Saudi Arabia). Or by /כך/ “kaph” if the dominant system was Hebrew (e. g. Israel).

One practical application of phonemic analysis is the creation of a phonemic notation to serve as a writing system (at least for those languages which do not already use such a system). A true alphabetic system uses one written mark (termed a grapheme, simply a letter) for each phoneme.¹⁶ The history of this kind of writing system is beyond the scope of this book, but for a number of reasons the traditional system used in the United States (and most English-speaking nations) is not usefully alphabetic. Consequently, English writing is not used by linguists except to represent meaning. Instead we will continue using phonemic versions of the phonetic notation graphs described in the prior chapter.

The last exercise for Chapter 3, 3–10, asks you to try your hand at coming up with phoneme symbols. Again we will refer to Zoque since you now have some idea of the phonemic system. Decide which symbols you will use for each set of allophones and state your reasons for your choice and then phonemically rewrite the selected items. In fact this is what you would be doing if you were trying to design an “alphabet” for a Zoque-speaking community.

Phoneme Sequences

There is one last topic which must be included when describing phonology and that is the system of patterning in phoneme combinations. Phonemes are arranged only in certain sequences to produce acceptable units. This means that some combinations are not permissible. At the beginning of this chapter we examined how many different total signals were possible given certain numbers of phonemes. For example, given a phonology of only 4 phonemes, and with the restriction that sequential repetitions of the same phoneme could not occur, there would be 36 possible 3-phoneme signals (all four phonemes times the other three times the other three— $4 \times 3 \times 3 = 36$). But actually there would be far fewer than 36 (if this hypothetical 4-phoneme language operated like real

¹⁶ Actually, this should be phrased as: one family of written marks per phoneme per written mode. All writing systems have different modes such as printing, script, capitalization. As long as the marks in one mode represent one phoneme each, then they satisfy the one mark/one phoneme relationship.

languages) since even without repetitions not all possible sequences could occur. Suppose the vowel phoneme could only occur in the middle position, and that the nasal phoneme could never be in the initial position. These two restrictions would bring the possible signal number down to 6 (the two non-nasal phonemes times the one vowel times all three consonant phonemes: $2 \times 1 \times 3 = 6$).

In English, for example, the velar nasal (/ŋ/) cannot occur in the initial position, while the glottal fricative (/h/) and bilabial and alveopalatal median resonants (/w/ and /y/) cannot occur in the final position. For fricative-plus-stop combinations, the manner of voicing must be the same—e.g. /s + b/ or /p + z/ are not permitted. These and other restrictions¹⁷ reduce the total potential number of English signal groups, but with 36 segmental phonemes this still permits more possible signals than are used in English. (There are a large number of permissible phoneme combinations which are not used—e.g. /mip/, /spaf/, and /nin/ just to give a few.)

The Syllable

The **syllable** is the significant unit in studying phoneme sequences. Sequencing patterns operate primarily within syllables. This can be seen in the English restriction on using voiced with voiceless stops and fricatives. The sequence */sba/ is not permitted, but only if it occurs within a syllable. The phrase “yes boss”, in fact contains this combination (/yesbas/), but it crosses syllable boundaries (the /s/ is at the end of the first syllable and the /b/ begins the second). Therefore a single syllable consisting of an /s/ followed by a /b/ and then a vowel, /sba/, would not be possible, while /spa/ “spa” is. The study of phoneme sequences is therefore dependent upon syllable identification, and, as we have seen above in chapter 2, while syllable peaks may be marked by a distinctive stress, syllable boundaries are not easy to establish. For this reason, most studies (such as the one by Whorf referred to below) are done on individual words which are clearly only one syllable. The information gained from this analysis can then be applied to identifying syllables in words of more than one syllable. Also, in words of several syllables some of the syllables are grammatical elements such as prefixes or suffixes which contribute to the words’ meaning and sentence function. These syllables usually do not occur separately as words, but they can be easily

¹⁷ A phonemic formula for a single English syllable, giving just the possible combinations, is presented in one of Benjamin Whorf’s classic articles, “Linguistics as an Exact Science”. The notation he uses is somewhat different from that presented here, but it is easily understandable. A modified version is presented in Figure 3–1.

identified in larger words since they can be added, taken away and replaced. Finally, the speakers may help identify syllables. Since the syllable is such an important element in the structure of speech utterances, speakers are aware of them (unlike phonemes and allophones) and can usually break large words into constituent syllables, or, when asked to say a multi-syllable word slowly, speak it syllable by syllable. In any event, much of syllable structure can be obtained, and with it much of the patterns of phoneme sequences.

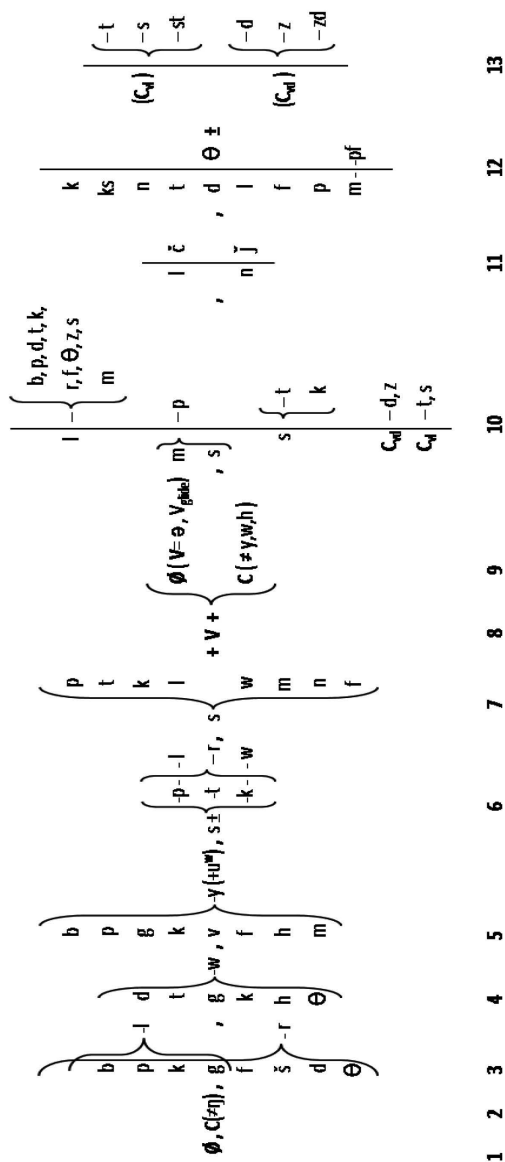
At this point we are able to provide an alternate definition for the terms **consonant** and **vowel**, which, as we learned in chapter 2, are the two kinds of phonemes in a language's phonology. In chapter 2 we gave a definition that was based on the degree of airstream modification. A syllable-based definition is as follows: a vowel is a phoneme which serves as a syllable peak, the part which carries the weight of stress and tone, while a consonant is a phoneme which serves as a syllable boundary, the part(s) which surround the peak either in the initial and/or the final positions. In effect, this means that the terms "vowel" and "consonant" must be understood according to each language's phonology—if a phoneme can serve to carry stress and tone as a syllable then it has to be understood as a vowel *in that phonological system*. For example, in Kiswahili (east Africa) nasal resonants often serve as vowels even though they were placed with the consonants in chapter 2.

Once the phonemes of a language are identified, the next step in describing the phonological system is to figure out how they are used in syllables. This is presented in two ways. The first is the sequencing of consonants and vowels which are permitted in the phonology. Some basic possibilities are CV, a syllable composed of an initial consonant followed by, and ended with, a vowel (e.g. the English "bee" /bi/). Another is simply a syllable composed of a vowel V (e.g. the English "a" /e/) and another is one with a vowel as the initial and ended with a consonant VC ("in" /in/). In some languages these are the only possibilities for syllable construction. English possesses a much more elaborate pattern. Not only syllables begun and ended with a consonant CVC (e.g. "cat" /kæt/), but in addition syllables begun and/or ended with two or more consonants strung together (called **consonant clusters**)—CCVC "slip" /slip/, CVCC "sect" /sekt/, CCVCC "slipped" /slipt/, and so forth up to even CCVCCCC ("glimpsed" /glimpst/). While English has all of the above types, and more, some languages have less. Polynesian languages, for example, do not permit any consonant clusters and so there would be no syllables containing CC (either as initial or final) as there are in English. Swahili does not permit a consonant in the final position and has only a few permissible consonant sequences. Actually, English is untypical among the world's languages for the number of consonant clusters it permits and a

word such as /glimpst/ would be considered improbable by speakers of many of the world's languages.

The second kind of patterning is the sequencing of specific phonemes or kinds of phonemes. For example, as we saw above, even though English permits an initial consonant in a syllable (CV), this can be any consonant except /ŋ/. Also, not all consonant clusters are permitted in English, the /-sb-/ sequence mentioned above, is not a permitted consonant sequence in English within a syllable. Figure 3–1 illustrates this second kind of patterning for English—it is adapted from a formula for constructing a single syllable word devised by the anthropological linguist Benjamin Whorf. You needn't worry about learning it, as an English speaker you already have it in your mental phonology. Note that it accounts for the basic syllable mentioned above: “bee” which is formed using part 1, any consonant except ŋ, and part 8, any vowel (parts 1 through 7 are alternatives and we used #1) and then that choice in part 9 which is no final consonant, or zero, as long as the vowel is ə or a glide. “Glimpst” would be constructed using part 3 (“g” + “l”), part 8 (“i”), part 10 (“m + p”), and finally part 13 (“s + t”), s would be a voiceless consonant, indicated as C, joined with a “t”).

This, then, is what the speaker must have for phonological competency: the articulatory characteristics of segmental phonemes and their allophonic variations, the suprasegmental phonemes (as well as the non-phonemic prosodic features of normal speech), and the patterns of phoneme sequences to form syllables and longer speech activity.



Notes: Overriding restriction on C combinations—CC cannot consist of same consonant

V = any vowel C = any consonant, C_{vd} = voiced consonant, C_v = voiceless consonant Ø = absence of a sound in this position

, = or (in a series of alternatives) + = and (in a series of required sounds ± = with or without the preceding sound ≠ = not including this (these) sound(s)

() = limiting condition for sound }_i { = included all in a combination (with preceding or with following sounds respectively)

| = combinations possible with sounds on both sides, unless linked by a double dash, --, which indicates only that particular combination, - (a single dash) = possible combination

Figure 3-1 Formula for Monosyllabic English Word
(After Whorf, "Linguistics as an Exact Science")

Terms to Know for Chapter Three

phoneme	allophone
phonology	allophonic consistency
distinctive features	assimilation
phonemic analysis	phonemic consistency
corpus	phonemic description
minimal pair	alphabetic writing system
minimal pair analysis	syllable
distributional analysis	syllable peak
suspicious pairs of sounds	consonant
complementary distribution	vowel syllable structure

CHAPTER FOUR

SEMANTICS: THE STRUCTURE OF MEANING

This chapter presents the linguistic study of meaning. However intricate the phonetic and phonemic patterning of speech may be, this is not the point of human communication—rather it is the transmission of meaningful messages among speakers and hearers. Speech sounds are used to construct signals having particular kinds of significance which are used under certain usual circumstances and in typical combinations. We shall begin with a discussion of the nature of semantic reference, especially the arbitrary association between signal and significance. Then, from the standpoint of the linguist studying the semantic system of another language, we shall examine how meaningful signals (called morphemes) are identified and the problems associated with this process. Finally, we shall look at some of the ways linguists organize the meanings in the semantic system of a language: the relation between “primary” and “secondary” reference, the extensions of meanings, types of meanings and relationships among meanings, and meaning domains.

The Nature of Meaning

The preceding chapter dealt with the nature of the sound signals used in human speech. We learned that speech consists of distinctive kinds of sounds (phonemes) put together in certain acceptable sequences (syllable structure). By learning the phonology of a language a speaker is able to produce appropriate sequences of speech sounds and thereby make acceptable speech signals. By itself, however, phonology is of no use since there is little that can be achieved or gained by phonology alone. If you were, say, a professional singer and needed to sing songs in another language, for example Japanese songs to a Japanese audience, it would be useful to master Japanese phonology even if you did not understand the meaning of the songs. However, this would undoubtedly become very frustrating and unsatisfactory for you. Humans acquire phonology so they can put together groups of speech sounds which have significance as messages. Semantics is the study of the system of signification employed by the speakers of a language in organizing their spoken messages.

Significance and the Precision of Message Content

By the term **significance**, I am referring to the appropriate use of speech signals as messages in communicational exchanges. A more common term is “meaning,” but unfortunately this may suggest something more precise, or perhaps more concrete, than is usually the case for the messages used among the members of a community. Messages associated with signals such as, in English, /²hə^wyə ³ðu:in¹/ “How ya doing” [noncommittal greeting], /³ó^w ¹bilyu^wbig²jejk²/ “Oh Bill, you big jerk” [affectionate chiding], /²wəts ³hæpinin¹/ “what’s happening” [generalized query], /²dæmit²/ “damn it” [mild expletive], /⁴gó^wfojit¹/ “go for it” [encouragement], /²a^w:j⁴áit/ “all right” [approval]¹, and so on, have acceptable significance as elements in speech interaction. These have “meaning” but not in the way that we would believe that the messages associated with the signals /krésint jéⁿnč/ “crescent wrench,” /jedwɪŋblækbjd/ “red-wing blackbird”, or /θjɪ¹á^wejzæn²fifti¹nminəts/ “three hours and fifteen minutes” have. Actually one of the essential characteristics of human messages (in fact, the single cause of the extraordinary cultural capabilities of our species) and one of the main reasons why semantics is the most challenging field of linguistics, is that human messages are of considerable complexity and elusiveness. “Crescent wrench,” for example, may seem to have a fairly definite meaning (which is why I used it as a contrast to terms such as “how ya doin”) but to an (older) mechanic even it is vague. It is actually a brand name for one of a number of adjustable wrenches that come in different sizes, shapes, manners of adjustment, and single or double combinations. In other words, just asking a mechanic for a “crescent wrench” may not “mean” too much. On the other hand many people have heard the term, and would be able to use it in a conversation (“Did you buy Dad a crescent wrench for Father’s Day?”), but couldn’t actually identify a crescent

¹ In order to discuss the topics in this chapter we must often distinguish between the message and its signal. Message content will be given in English and presented in a spelled form but enclosed in square brackets []. Signals will be given in phonemic notation, enclosed by slant lines / /, and followed, for clarification (although if the student has diligently memorized chapter 3 this shouldn’t be necessary), by a spelled form enclosed in quotation marks “ ”. For example, [a small spherical object made of a cork center wrapped in cording and covered with leather] /bé’sbal/ “baseball.” If the signal is particularly lengthy, or if it would aid in ease of sentence structure, the phonemic representation of the signal may be omitted and only the spelled form given. However, the author firmly believes that it is important to give as accurate a phonemic representation as possible, including stress (superscript numbers) of meaningful units.

wrench in a tool box containing a variety of tools. For these people, the term “means” something, but perhaps only as [a member of a set of machine-type accessories that are suitable as Father’s Day gifts]. And certainly, for many the term means nothing. Furthermore, for a parent who dreams of a teen-aged child becoming a lawyer, the child’s request for a “crescent wrench” as a birthday present may well mean [disappointment and the dashing of one’s hopes].

The Range of Significance for Any Speech Event

To illustrate the range of significance which a single utterance can have, consider the common phrase, “please pass the salt.” This could mean:

- [a request to pass the salt, located in a small container, during a meal];
- [a request to pass a piece of salt, such as rock salt, among a group of minerals];
- [a statement that the food is tasteless, or insufficiently salted];
- [an accusation that the cook did not put in sufficient salt—as was expected];
- [an announcement that the speaker intends to go off a low-sodium diet];
- [an announcement that the speaker has recovered his/her sense of taste];
- [a mocking reference to another diner’s insistence upon formal table manner, i.e. instead of just saying “Gimme the salt”];
- [an attempt to start a conversation with another diner, one closer to the salt container];
- [a warning to certain, probably young, diners, not to throw the salt, or otherwise play with it]; and so on.

What all this should suggest is that meaning is rarely precise, and concrete, as a message such as [profanity] /dæmit/ “damn it” illustrates (to say nothing of the meanings of such items as [space], [electron], [time], [democracy]). In fact no one should approach the study of semantics with the belief that meanings must be precise. They are likely to end up writing a diatribe against their informants’ sloppy thinking because they assume incorrectly that a (supposedly) precise non-language system like logic or mathematics is the ideal against which the message content of human speech communication can be assessed.

Phonology is based in human physiology and so has universal constraints on it according to the principles of articulation and the physical

properties of sound signals. Consequently, as we have seen (chapter 2), a single descriptive system (articulatory phonetics) can be accurately used to describe any human's speech. On the other hand, semantics is primarily based on custom and tradition. Each culture is its own system of meaning. Universal semantic characteristics are not readily apparent and therefore there is no single descriptive system which can be used to describe different semantic systems equally well. Instead, the semantic system of a language being studied is described in terms of the semantic system of another language—typically the linguist's own language, typically one of the world's dominant languages such as English.

The Arbitrary Nature of Message Content

Since semantic systems are customary and traditional rather than being based in physiology, the association between message and signal is essentially arbitrary. There is no inherent relationship between any particular message and the sound signal which it is associated with except by the ongoing convention of the community of speakers. For example, a speaker must learn the significance of the U. S. English signal / bé'sbal /. Nothing in the signal itself in any way contains this significance (or any significance at all). It is certainly true that many American English speakers will come to believe (quite firmly) that the signal / bé'sbal / must necessarily mean [the game played with a hard ball between two teams of nine players who alternate as batters and fielders according to a particular set of procedural rules, a pennant race, sliding hard into second to break up the potential double play....]. They would assume that /bé'sbal/ can not mean anything else and no other signal could have quite the same significance. To the person (in any community of speakers) who learns the significance of common and important signals it is as if there is no difference between a signal and its significance. This belief, however, is not a form of proof that signals must mean certain things—it is rather a testament to the power and effectiveness of cultural learning. In Spanish speaking communities, however, the signal is /be'z ból/, in Japan it is /bé'zubôru /². Now although you might object that these are essentially still the same signal and that this shows that signals and meaning must correspond, consider that the signals have been modified and in fact this proves that signals are not absolute determiners of meaning. If there were a necessary relationship then a signal could not change at all and still be associated with the same meaning.

² The original term was /yakyu/ “fieldball”. /bé'zuboru/ is an old English loan, however some of my Japanese students have told me that a form more similar to the American /bé'sbal/ is preferred.

Approaching this question from another angle, we can also see that even if the signal does not change, its meaning may. The signal / bé'sbal /, for example, even if it was not modified, would not have exactly the same significance in America, Cuba (as one of the Caribbean countries where it is played), or Japan. Those parts of the total significance of / bé'sbal / in America including [the national pastime, hot dogs and beer, and the verbal abuse of players and umpires by fans] would not be carried along to other cultural settings.³ So signal form and meaning content are not inherently bonded to each other. Anthropologists refer to linguistic signals and the meanings associated with them as **symbols**. The only way to figure out the significance of symbols is to learn them. Nothing in the form of the symbol indicates what its significance would be and nothing in the content of a particular meaning indicates what its associated signal would be. This flexibility is what gives human communication its extraordinary message carrying capability. (Otherwise we would have to communicate by some sort of verbal charade—imitating the sound of something in order to refer to it.)

Openness and Productivity

Phonology provides the means for a very large number of different signals and these can be used for a potentially unlimited number of messages. Moreover, signals and messages can be modified as required to create new forms and content independently of one another. This is why human communication is both **open** and **productive**. Language displays openness because there is no limit to the number and kinds of messages which may be used by a community of speakers other than that provided by the practical boundaries of their culture. Arctic hunters (e.g. Inuit caribou or seal hunters) do not need to converse about a tropical forest three-toed sloth, and tropical farmers, similarly, do not require conversation about polar ice conditions. Still, the range of what each community may talk about is very extensive. It includes a wide range of social relations, human emotions, desires and capabilities through their life cycle, work and technology, decisions and social control, and the many aspects of the supernatural and its relationships with humanity, to name a

³An excellent example of the change in a signal's significance as it is used by a different culture is depicted in the film *Trobriand Cricket* (produced by Jerry W. Leach). The British signal /kriket/ is applied by the members of the Trobriand Islands to a somewhat similar game but one with considerable different public, competitive, and performance significance. Just to take one example, Trobriand cricket players, in stark contrast to their British counterparts, use the occasion of a match to ornament themselves in extravagant fashion and engage in choreographed dancing and chanting at regular occasions throughout the match.

few of those dealing with humans themselves. Moreover, the members of a community can always talk about something new, developing new signals and meanings in the process. And related to this is the human capability to use existing messages in non-usual combinations and contexts thereby creating ever differing total messages.⁴ This ability to rearrange meaningful units and thereby create new messages is called productivity. This is a crucial aspect of the creative capabilities of human language. It is the basis for innovation and cultural adaptation, to say nothing of its importance in verbal play and aesthetic performances. In the absence of openness and productivity it is doubtful that humanity could have occurred.

Morphemic Analysis

Minimal Components of Message Content

We will begin our study of semantics by seeing how we can identify those signal groups which have significance in a corpus. Our strategy will be to try to isolate the smallest such units because an important assumption of semantic analysis is that the significance of total utterances lies in the combined significance of their individual components. The utterance “Will you stop bothering the cat?” certainly seems to have a total significance because it often has (what appears to be) a single communicative result: The addressee stops bothering the cat. Yet this significance is the result of its parts. This can be easily demonstrated by replacing parts, (“Will you stop beating the cat?”), removing parts (“Stop bothering the cat!”), or rearranging parts (“Will the cat stop bothering you?”) and thereby changing the significance of the utterance.⁵

⁴ There are some practical limits to productivity in regard to individual speakers, however. Productivity is a potential for innovation and creativity. In other words, we do not just repeat the same messages over and over again (although a surprising amount of conversation is repetitive). But an individual who produced unusual message combinations in every utterance would quickly be ignored or treated as insane or a fool.

⁵ In fact multi-unit utterances often have a significance which is not the same as the sum of their parts. Idioms are one example of these—“son of a bitch,” for example, does not have the significance of the sum of its four parts. The idiomatic nature of utterances (as important an aspect of any semantic system as they are—in many respects they operate as if they were single signals) will not be discussed here. However, even for such utterances as idioms, speakers can recognize that they can, indeed, be treated as ordinary combinations and these parts can be modified with subsequent corresponding effect on the significance. “Son of a

Identifying Morphemes

Identifying minimal signal groups which are associated with meaning can be a straight-forward process of determining those phoneme sequences which co-occur with the same meaning. The actual process is more challenging than it sounds. We will assume that we have enough of the phonology to transcribe all the speech items in a corpus into phonemic notation. (Semantic analysis, by the way, doesn't have to wait until all details about the phonology have been resolved.) While a phonemic representation is not absolutely necessary, it is much more efficient because allophonic details of speech behavior have no effect on meaning. A phonemic corpus thus allows the linguist to just examine speech sound differences or identities which are actually used by the speakers to signal differences or identities of meaning—e.g. phonemes.

Corpus and Gloss

The corpus will include a **gloss** for each item. The gloss is an item in the linguist's language that corresponds to the meaning of the item in the target language. The significance of the gloss will (hopefully) be similar to that of the item. The simplest corpus will be one where each item is a single signal having a particular significance corresponding fairly well to the significance of the gloss. Putting aside for the moment the question of the appropriateness of the gloss, let us look at the situation where the speech item is not a single signal. (This is much more common, partly because humans usually do not employ utterances containing only one signal, and partly because in many languages even the smallest possible utterance contains two or more signals.)

Co-occurrence of Phoneme Sequence and Gloss

Here is a sample corpus of three items from Kiswahili (spoken throughout Eastern Africa):

- | | | |
|-----|-----------|--------------------|
| 4.1 | /anačéka/ | “he (is) laughing” |
| 4.2 | /alipíka/ | “he cooked” |
| 4.3 | /ameséma/ | “he has said” |

Because the glosses consist of several meanings, we will assume that items 4.1—4.3 also consist of more than one meaningful part (this will not

birch” might be used as a joking reference in regard to a forester or woodsman, for example.

always be a useful assumption). In order to identify the small parts of the items we must find recurring phoneme sequences which co-occur with the same meaning (for which we will, by necessity for now, use the gloss). Examining the three items we find that only the initial /a-/ phoneme and the final /a/ phoneme occur in all three items. (Actually, the stress on the next to last vowel also recurs but for simplicity we will assume, rightly as it would turn out, that this is not a meaningful signal.) Examining the gloss we find that only “he” occurs in all three glosses. As an initial hypothesis we could say that the phoneme sequence /a---a/ is a meaningful element in each item and means “he.” (Don’t be put off by the idea that these two phonemes are not next to each other. A discontinuous phoneme sequence is, in fact, one possible kind of signal.) The requirement of co-occurrence, however, is that the phoneme sequence in the speech item must occur only with the gloss item and, vice versa, the gloss item must only occur with the phoneme sequence. (This “requirement” will turn out to be rather flexible, but it will serve for now.) Suppose we added the following items to our corpus:

- | | |
|--------------------|----------------------|
| 4.4 /ninačéka/ | “I am [is] laughing” |
| 4.5 /alijaribu/ | “he tried” |

In item 4.4 the final /a/ remains but the “he” gloss is absent along with the initial /a/. This would seem to support the conclusion that the initial /a/ and not the final /a/ corresponds to the gloss “he”. Item 4.5 verifies this conclusion. The final /a/ is absent but both the initial /a / and the “he” gloss are present. The final /a/ may or may not be a meaningful unit of speech, but it has nothing to do with “he”, instead it seems clear that only the initial /a/ can be identified as the meaningful unit of speech associated with the English gloss “he” (the fact that it is a single phoneme is of no concern since it is still a phonemic unit, albeit a minimal one)⁶. **It is also important to state that when we mentioned the requirement of co-occurrence, what was meant was both co-presence and co-absence.**

We will call /a/ a **morpheme**. This is the smallest phonemic unit having significance for a group of speakers. This significance is based on its contribution to the message of an utterance. (Do not assume that morphemes are single phonemes—some are, but most consist of sequences of phonemes.) Thus a Kiswahili speaker who wishes to refer to the action of someone not in the immediate group will use the /a/ morpheme in

⁶ Identifying morphemes is often no more than repeating this process. However, the task becomes easier as more and more morphemes are identified. Knowing that /a-/ is “he” allows the linguist to use item 4.1 and a form like /tunačéka/ “we are laughing” to tentatively decide that /tu-/ is “we” in the gloss even though there is only one instance of this form.

constructing an utterance (e.g. / ĵana bwana mutiso aličeka / having the English gloss: “Yesterday Mr. Mutiso (he) laughed”).

Morphemes or Words

Morphemes, then, are the basic building blocks of meaningful utterances. However, they may not always occur by themselves as complete utterances, i.e. as separate items in a corpus. A Kiswahili speaker will not, for example, say /a/ as a complete utterance meaning “he”—for example as a response to the question /nāni ¹aličeka³/ “Who laughed?” (Actually “he” would not usually appear in English as a response to a question such as this, rather “him” would be the appropriate response, or “he did.” However, we would probably provide “he” to a question such as, “What is the word for a first person singular male subject?”) It is also likely that this speaker might not be able to tell a linguist that /a/ means “he” for the same reason that an English speaker might not be able to tell a Kiswahili linguist that /s/ means [more than one, plural] when it occurs in an item like “cats”—neither speaker is accustomed to using the morpheme in isolation. Still, each speaker ‘knows’ their significance because they use the morpheme in appropriate ways, the Kiswahili speaker can replace /a/ with /ni/ and change “he is laughing” to “I am laughing” and the English speaker can leave off the /s/ and change from [plural cat] to [singular cat].

The smallest possible utterance in a language is called a **word**. In English many words are also single morphemes (such as /kæt/ “cat”) while in Kiswahili most words are combinations of morphemes (such as / a + na + ček + a/ “he is laughing”). Speakers are familiar with words, even if they may not have an explicit knowledge of their constituent morphemes, and these are the speech items which linguists will elicit from their informants for their corpus. Utterances consisting of more than one word are called **sentences**, if they are felt by their speakers to be complete utterances, or **phrases** if they are not. (In chapter 5 we shall use intonation for the definition of a sentence.) Multi-sentence utterances, usually involving more than one speaker, are discourses, conversations, or “texts” (especially if they are stylized forms such as poems, myths, folktales)

At this point you should try your hand at identifying morphemes from a corpus. Look at the example at the beginning of the Exercises and make sure you understand it. The first three exercises for Chapter 4 are relatively straightforward matching of repeated phoneme strings in the items to repeated elements in the glosses and so you should be able to figure out the morphemes in each corpus. Exercise 4–3 is a longer corpus and has more morphemes to identify (but the same principles of identification

operate), and also has a small challenge in morpheme identification that is possibly covered below in the discussion of homophones.

Glosses and the Question of Morpheme Significance

It might seem that the presence of glosses solves the problem of figuring out the meaning of the items in the corpus but this is clearly not the case. There still remains the problem of determining the smallest units with meaning, and items may actually be composed of several smaller units and this composition may not be obvious from the gloss. It is important to keep in mind that the gloss is definitely not an accurate definition of the significance of an item, even if an item is not divisible into further units. The gloss is an approximation by the linguist. This is usually all that is needed for phonological analysis since the only semantic relationship it requires is same—different. Obviously semantic analysis requires much more than this relationship (although same—different is one of the basic semantic relationships). But there are many limitations to a better approximation being provided by the linguist for the gloss.

First, it is extremely unlikely that the semantic significance of any item in the language being studied will be matched by the semantic significance of an item in the linguist's own language (or any other language) although there would certainly have to be some overlap. The Nuer (Eastern Africa) word “ghok”⁷ would probably be given the gloss “cattle” by an English-speaking linguist. And, in fact, while this is not completely inappropriate (as would be the case if the linguist had instead given a gloss of, say, “horse”) it is not completely accurate either. The Nuer animal is the Asian long-horned humped version of the *bos* species, not the dairy farm Guernsey probably thought of by most Americans. More importantly, to the Nuer “ghok” has been described as “meaning”: wealth and prestige (suggesting a gloss of “Jaguar sports car” or “Cadillac”), marriage and fertility since they are used for bridewealth (suggesting vaguely a gloss of “diamond engagement ring”), fond loving care (suggesting a gloss of “pet”); and even a means of religious sacrifice (suggesting a gloss of “offering”). In other words, the single English word “cattle” is not a complete gloss since the significance Nuer associate with “ghok” is not the same as that which most Americans associate with “cattle.”

⁷ A society made famous in the writings of Sir E. E. Evans-Pritchard. They live in the upper Nile Valley in Sudan and Ethiopia. Only a spelled version of the Nuer word for cattle is given here since different sources give different representations and the author has no access to a native Nuer speaker. In our phonetic notation, however, it seems likely that it is /yak/.

This need not be an impediment to an English speaker trying to understand Nuer since anything expressible by one community of speakers is ultimately expressible by another. The problem in linguistics is that we are trying to start with fairly short utterances (if not single items) in the language being studied so as to isolate units with meaning. It would be awkward to have a corpus where every item such as “ghok” is associated with a correspondingly lengthy English gloss. In this respect the linguist doing semantic analysis may become more like a translator who searches for the most efficient (i.e. shortest) match as a gloss rather than the most completely accurate.

The reader should at least realize that the common practice of presenting short glosses for items in another language (in vocabulary lists for language courses, handbooks on other languages for travelers, as well as in most linguistics exercises) can be misleading. However it is certainly true that as long as the linguist understands that a short gloss is very likely only a beginning approximation, and is simply using this as a way of establishing a list of items which have (some) meaning, it is permissible.

One solution is for the linguist to create a “dictionary” or lexicon listing the many references of each item and then simply identify these by number—e.g., for “ghok”[cattle₁, cattle₂. . .]. This may make for simpler glosses, but it still begs the question of obtaining these meanings in the first place.

Morpheme Significance Not Accurately Expressed in another Language

Ideally the linguist would have been able to observe actual speech interactions among community members and tie the occurrence of particular phoneme sequences to the conditions and results of their usage. This, not surprisingly, is exceedingly difficult, and, of course, presupposes that the linguist is already accomplished as a user of the language. More commonly, the linguist works with a bilingual informant. This person typically has had some formal schooling with the linguist’s language as the medium of instruction and has learned the kind of item for item correspondence where single words in the informant’s language are associated with single words in the linguist’s language. Thus a Nuer informant with several years of primary school will be asked by the linguist “What is the word for cattle?”, and will respond, “It is ghok!” This Nuer has undoubtedly not lived in the linguist’s country and doesn’t know if this is an appropriate semantic match (i.e. whether Americans attach the same significance to “cattle” as Nuer do to “ghok”)—he or she has simply learned it as true by the weighty authority of a teacher or, even more authoritative, by a written Nuer-English vocabulary list.

Glosses Only as Approximations of Morpheme Significance

Even if the informant is not bilingual, and the linguist is simply pointing or acting to elicit speech items (i.e. a linguist working with a Nuer informant pointing to a cow) there is no guarantee that the gloss will be close to the item's significance. The linguist might still elicit "ghok" and write "cow" or "cattle" as the gloss. More likely, the informant will provide a speech item different from what the linguist believes is being elicited. For example, the linguist points at a cow wanting a generic term but the Nuer informant provides the name of that particular cow (i.e. "luthrial"), or perhaps the term for a cow with those particular markings and shape of horns, or the term for which stage of life the cow is in (e.g. "ruath"—specifying a heifer); and the linguist writes down "cow."

In these cases, the speech item's significance is apparently more extensive than that of the gloss. This can go the other way as well. The /a/ morpheme in Kiswahili, for example, was given the English gloss of "he". But this English signal also carries the **embedded meaning** of [male] which the Kiswahili signal does not. A more useful English gloss for /a/ would be [third person, singular]. Of course, it is possible to achieve fairly good correspondence in translation, given any language's flexibility and capacity for graded comparisons such as [it is a little like X, but also includes Y]. Consequently, an English user can come to appreciate the significance of the Nuer signal "ghok" when the gloss "cow" is expanded with further associations. The important point is that the glosses of a corpus being studied by the linguist are not the same as the semantic significance of the items. The gloss is only a (hopefully useful) start to semantic analysis and also grammatical analysis, as we shall see, but is directly applicable only to phonemic analysis for which it is only necessary to know that two phonetically represented items have the same or different gloss.

Now try to do Exercises 4–4, 5, and 6. In each of these problems matching the gloss to the repeating phoneme strings will be a moderate challenge. Exercises 4–7 and 8 will stretch your ability to figure out embedded meanings in the gloss.

Description of Semantic Systems

Componential Analysis

Our study of semantics will be based on the significance of morphemes and their relationships. We will begin by imagining that the semantic system is like a dictionary. This will turn out to not be a good analogy, but it has certain useful features that illustrate some important aspects of

semantic systems. First is the concept of a listing of morphemes (the phonemic signal) alongside some description of its significance. Just as you can turn to an entry in a dictionary and see its meaning, members of a community can ask one another to state the meaning of a speech signal. One task of semantic analysis is to do such a listing—i.e. each morpheme that is identified is attached to a statement of its meaning. One branch of semantics developed in anthropology to define meanings is **componential analysis** (also called formal semantic analysis, or ethnoscience). Its objective is to provide quite precise statements of morpheme (or word) meanings by systematically examining how the morpheme was used. The basis of this approach is that morpheme meanings are not limited to just individual things or events (excepting perhaps semantic references such as personal names). Instead the semantic reference of each morpheme refers to a category of things or events. The members of each category referred to by a morpheme (or word) must share a set of distinctive attributes (i.e. the necessary components of its definition). The speakers of a language learn to apply morphemes by learning the attributes of each category. In other words, I will not be viewed as an appropriate speaker if I use a morpheme in reference to something when it lacks any of the appropriate attributes which are supposed to be present for that morpheme.⁸ For example the morpheme /buk/ “book” refers to a category of objects whose attributes include a number of sheets of paper (pages) bound together in a heavier binding which wraps the pages on three sides, allowing these pages to be turned over and kept together. The object you are now reading from is one member of this category, but there are a very large number of others. Many features associated with the members of this category, however, are not significant attributes. For example, size, color, content, and firmness of binding material, are each features of “books” but do not really play a role as necessary attributes (distinctive components) in deciding whether a particular object is or is not a “book”. The fact that something is “red” in color, by itself, does not make it a member of the “book” category. But in order to use the morpheme /buk/ appropriately⁹ a speaker must be able to recognize those objects which do have the

⁸ This issue is illustrated by Charles Dodgson (Lewis Carroll) in his *Through the Looking Glass* when he has Humpty Dumpty saying “There’s glory for you” to Alice and when she asks what he means by that, he replies, “I meant ‘there’s a nice knock-down argument for you!’” Alice protests that “glory” doesn’t mean that and Humpty Dumpty replies (and this is where he becomes, at least to those on this side of the Looking Glass, an inappropriate, even insane, speaker), “when I use a word, it means just what I choose it to mean—neither more nor less!” (Dodgson was actually poking fun at a philosophical debate about words having essential meanings.)

⁹ At least with regard to its core meaning (see below).

distinctive components of being a “book” from those which do not. If this were not the case human cultural interaction would be impossible. (The reader should keep in mind that given the productive nature of human communication, and our ability for metaphor and analogy—see the reference to “pigeon” below—the context in which a morpheme is used plays a significant role in deciding whether or not it is being used appropriately (for example in Shakespeare’s *The Taming of the Shrew*, Act IV Scene V, when Petruchio gets Katherina to say that the sun is the moon and an elderly man is a young maiden).

Componential analysis attempts to determine the set of necessary attributes (distinctive components) associated with each morpheme reference. This is a feasible approach for a number of morphemes, especially for those whose significance is to refer to a physical object or event. The analysis proceeds by recording all the features of those things to which a morpheme is applied. When a suitably large number of morpheme applications have been recorded, the analyst determines which features are present in all applications and these are then identified as the semantic category’s attributes. These attributes cannot be obtained by just asking informants for them because they have not been learned directly, but rather indirectly over a period of time by trial and error and observation from early childhood on. One of the contributions of this type of semantic analysis is that it should yield information about the actual workings of a semantic system which is more insightful (or “deeper” if you will) than can be provided by the speakers themselves. But this kind of method does have the very real advantage of being based on usage by speakers and it thereby fits into the anthropological linguist’s orientation of studying what ordinary humans actually do when they speak rather than some sort of idealized projection of what humans are supposed to do.

The perceptive reader might have also noticed that componential analysis is very much like phonemic analysis (chapter 3). In fact, it is a derivation of it. Phonemic analysis examines many instances of similar, but non-contrastive, sounds to determine what phonetic attributes they share as a way of defining the phoneme category to which they belong. It also assumes that phonemes will not be definable by speakers because they have been learned indirectly, early in life and are now part of speakers’ unconscious knowledge.

Limitations of Componential Analysis

Componential analysis is certainly consistent with a “scientific” approach to language since it promises a rigorous set of semantic definitions and this is one requirement of science. Unfortunately, the bulk of morphemes in a language do not, in fact, have neatly precise

significance for their speakers (as was pointed out at the beginning of this chapter). There are terms which require exact application, from technology to ritual and kinship, but these are in the minority. Remember, we are interested in actual usage, not ideal usage or just usage by specialists. The underlined words in the following English phrases are examples of common morphemes lacking precise significance: “I have a cold,” “it is a little further,” “let’s go out,” “it’s so political,” and so on (to say nothing of the significance of such morphemes as “have,” “it,” “a,” “go,” and “so”).

Core and Extensional Meanings

Denotation and Connotation

One traditional aspect of morpheme meaning is the distinction between **denotative** and **connotative** meanings. Denotation refers to the supposedly standard, or commonly expected, significance of a morpheme, while connotation would refer to all the extensions of the denotative meaning. For example, “pigeon” has as its denotative meaning a type of bird, but it can also refer to a gullible person who is being used or manipulated by someone else (e.g. a con man or swindler). This extensional aspect of morpheme meaning is present for all morphemes (to a greater or lesser degree) and in all languages. This is often viewed as if denotation was a basic (even correct) meaning, and then connotation a group of subsidiary (or even incorrect) meanings. The basic meaning would therefore be labeled “standard” or “formal” and the subsidiary meanings labeled “colloquial,” “slang,” or “dialect.” This, in fact, is the case for our example of “pigeon.” The meaning “type of bird” is given in my (Merriam-Webster) dictionary as the standard meaning and the meaning of a gullible person is listed as slang¹⁰.

The Question of the “Superiority” of Denotative Meanings

Descriptive linguistics strives to study the ordinary speech behavior of a community. If the members of a community in fact believe that one meaning of a morpheme is superior to another, and that the use of the less desirable meaning marks the user as possessing less desirable social qualities, then the linguist may be justified in applying a correct/standard—incorrect/non-standard (slang) label to the semantic relationship among the morpheme’s meanings. This kind of labeling is

¹⁰ Stuart Flexnor, *I Hear America Talking* (Simon & Schuster, 1976, p. 282) states of “pigeon” that “...it also meant a dupe by 1590 ...”

typical of dictionaries, and it is usually done as a part of political domination by that elite group which controls the writing of dictionaries for a supra-community entity such as a nation. In fact, this kind of assessment of semantic relationships (indeed of any aspect of language) is the subject of the subfield of sociolinguistics which studies the social (including political) dimensions of language. We have already referred to this approach to language in the first chapter under the term sociolinguistics.

Here, in contrast, we will treat the question of extensional meanings as a matter of frequency. A denotative meaning will be that meaning (or meanings) which are typically used and therefore expected by hearers, and connotative meanings are those which are used less often.¹¹ Thus, while accepting the concept of a core meaning (i.e. denotation) and extensions of it (i.e. connotations), descriptive linguistics does not characterize the relationship between these meanings as superior to inferior, just typical to non-typical. Extensions must, however, share some semantic properties with the core meaning. Otherwise there is no extension of semantic relationship but rather two different morphemes which happen to have the same phonetic shape.

The Problem of Homophones (not homonyms)

The [bird] and [gullible person] meanings of “pigeon” are considered to be sufficiently related in terms of their meaning, although perhaps merely just long associated with each other, that they are treated as core [bird] and extension [dupe]. But another “pigeon” spelled “pidgin” (we are, remember, dealing in speech not spelling and “pidgin” has the same phonemic shape as “pigeon”: /piʃin/) is treated as a separate morpheme with the meaning [a rudimentary blend of two or more languages used for restricted social intercourse among two different speech communities; a *lingua franca*]. Morphemes with the same phonemic (spoken) shape but with meanings considered to be unrelated are called **homophones** (homonyms refer only to spelled identities).

The distinction between: A) grouping several meanings together as core and extended meanings of the same morpheme, or B) distinguishing these from each other which, although associated with the same signal, are

¹¹ The difference in frequency of use among morpheme meanings is undoubtedly a function of a difference among contexts of usage. I.e. “pigeon” [bird] used most frequently in an urban park by the people there to enjoy the park while “pigeon” [gullible person] used less frequently and only by those in the park to swindle. Studying the relation between kinds of usage and the context of speaking is of particular interest to sociolinguists.

treated as really belonging to different, homophonous, morphemes is never an easy, or obvious one. It is only possible when done by the speakers themselves, not the linguist. In fact, it, among many other intellectual tasks, is one of the responsibilities of scholars as they set out to establish or rework dictionaries. Homophones exist in all languages but are of no direct interest to semantic analysis (although, as the reader should appreciate, their presence presents problems to linguists as they proceed with their analysis.) Rather, one of the important steps in semantic analysis is to examine the extension of meanings when they are considered similar enough to be still grouped within the same morpheme. Extension of meanings is an ongoing feature of human communication as speakers use their considerable creative abilities to apply morphemes in slightly (or radically) new ways to the changeable physical and social world in which they live. Morphemes are typically used over and over again with the same significance. But the human capability for analogy coupled with the wide range of sentence or phrase settings in which morphemes can occur gives each use of a morpheme the potential for extension of its meaning. Thus:

- 4.6 “That bird is a pigeon”
- 4.7 “He walks pigeon-toed”
- 4.8 “Pigeon Bill always feeds the birds”
- 4.9 “Nice going, pigeon brain”
- 4.10 “Statues hate pigeons”
- 4.11 “He made pigeon-sounds of delight”
- 4.12 “She was a strutting pigeon in her pride”

and even Gertrude Stein’s poetic references to a miraculous vision of the Holy Ghost from her libretto to the opera “Four Saints in Three Acts”

- 4.13 “Pigeons large pigeons on the shorter longer yellow
grass alas pigeons on the grass. If they were not
pigeons what were they.”

all elicit somewhat different meanings of “pigeon” according to the differing morpheme environments in each phrase. But for as much opportunity as there is for metaphorical applications (and this is considerable) there are limits. There are boundaries around “pigeon” and therefore phrases such as 4.14—16 will receive a response of “Huh? What do you mean?”

- 4.14 “Will you pigeon the salt please”
- 4.15 “My pigeon ‘tis of thee...”
- 4.16 “Hand me that pigeon head screwdriver”

Exercise 4–8 gives you an opportunity to examine the manner in which a dictionary handles some common English words semantically. You can assess the appropriateness of what the dictionary editors came up with based upon your own usage.

Semantic Categories and Relationships

What we have seen so far is that one of the tasks of a semantic analysis is to describe the areas of meaning of the language's morphemes. Even though each language has its own set of morphemes used in ordinary speech interaction (numbering in the ten to twenty thousand range) and there is a uniqueness in each language's set of morpheme meanings, linguists can assume that there are certain general categories of meanings present in all languages. This area of linguistic study, language/semantic universals, has as one of its goals the discovery of the general properties of how the mind operates. All we can do here is to tentatively provide some of the basic ideas in the study of universal semantic properties. For example, **nouns** may be defined as morphemes or words whose semantic significance is to refer to some entity about which something else can be related. Our term "pigeon" can be identified as a noun and it can be related to an activity ("pigeons fly"), a quality ("fat pigeons"), a comparison ("like a pigeon does"), a location ("in the pigeon") and so on. **Verbs** possess the semantic significance of activity either done by or to a noun ("The pigeon lands on the statue", "The statue sits on the pigeon"). An **adjective** possesses the semantic significance of a quality as an attribute of a noun ("The dead pigeon is on the ground"), and an **adverb** similarly represents attributes of a verb ("The pigeon carefully preens").

It is necessary that these semantic types (meanings) be tied to some corresponding spoken usage of a morpheme or morpheme arrangement (which will be dealt with in the following chapter). For example gender can be discussed in any language, but to be included as a semantic category in the description of a language it must be associated with a morpheme. For example, in Spanish the noun suffix morpheme /o/ is [masculine gender] and /a/ is [feminine gender]. Kiswahili has no such morpheme (as part of the morphemes which combine to create nouns) and so gender would not be included in a description of Kiswahili as a semantic category (whether as a noun, adjective, verb). There are many other types of semantic categories which may be universals, such as locatives, articles, demonstratives, verb tenses, noun function (subject—object). These are all necessarily general categories. Some are present but not too useful for some languages, absent in others, and some are crucial in others. Each one should not be assumed to be present in every language, but they can be starting points in a semantic study. Most importantly, these

must not be confused with the particular ways they operate in any one language, **especially English**. Just because English is a language of economic and political power in today's world should not confuse one into thinking that its semantic categories are therefore universally important and necessary for all languages. (This last potential for confusion may be avoided if we use more obscure labels such as “nominals” instead of “nouns”, “predicates” for “verbs”, “nominal modifiers” for “adjectives,” “verbal modifiers” for adverbs, and so forth.)

Of greater interest, perhaps, than types of semantic categories, are types of semantic relationships. Sub-categories of nouns such as **subject** and **object** indicate such relationships as initiator (or agent) - recipient (or patient). For example, “Tom greets the president” or “Jane started the motor.” Also involved in this relationship is the verb sub-category of **transitivity** which refers to an activity which has an effect on something (“Mary kisses James”) as opposed to **intransitivity** (“James faints”) in which the effect, if any, is with the initiator of the action.

Another category is **voice** (not to be confused with the articulatory term “voicing”) which indicates the direction of action between subject and object. In an active voice the subject's place in the phrase is the initiator of action, in a passive voice the subject's place is the recipient of the action, for example, “the president was greeted by Tom”. There are many other potential verbal relationships: a **conditional** category where one action or event is dependent upon another (“If she arrives today, pay her”), **reciprocal** where both the initiator and recipient of an action are the same (“He cut himself”), **causative** where an action on an object is specifically caused by a subject (“Jim broke the vase”), and **aspect** which indicates the distribution of an action in time, such as an action which occurs over and over again (“The waves pound the rocks”), to name only a few. In adjectives there are semantic relationship sub-categories of **comparison** (“the smaller girl”), **exaggeration** (“ultra-hero”), **relational** in which a noun is modified by a phrasal structure (“the man who would be king”), and **possessive** (“Tom's book”), to list only a few. Other such relationships involving semantic categories are **conjunctive** (“John and Mariana”), **copulative** in which the subject is provided an identity with something else (“Kim is a freshman”), **locative** providing location and direction (“in the soup”, “down the street”), **negation** (“She did not say”), and several **associative** categories which provide one kind of thing (e.g. a person) with the qualities of another kind of thing (e.g. an activity) such as the noun “runner” from the verb “run”, or simply provide a linkage among different elements in a phrase (in Kiswahili /kitábu kidógo kíle/ “That small book” where each modifier takes the same class prefix as the noun). These latter categories are often referred to as having a “grammatical” meaning.

A linguist has the enormous task of describing a language's system of meaning categories and the many kinds of relationships among them. The brief discussion above only provides some of the categories and relationships which linguists have found in many of the world's languages but is by no means exhaustive. Furthermore, it must be emphasized again that although the examples above were from English (with the one exception from Kiswahili), the linguist cannot rely on a gloss which matches the English format or words. For example, the presence of "in" in an English gloss is no indication of a locative semantic category as the possible gloss "in the mood" demonstrates.

Semantic Domains

A Dictionary Model

There are many other semantic relationships. Morpheme meanings occur in clusters according to their overlapping concern for a common area of human interest. We will call these clusters **domains**. Let me take a brief digression to explain the importance of the domain concept. We mentioned earlier that a dictionary had some use as a model for a semantic system, at least for that part of the system which dealt with morphemes and their significance. We also saw that there were limitations to this model in regard to its description of meaning extensions.¹² However a much more serious limitation lies in the manner in which a dictionary is organized. As a by-product of its system of writing, an English dictionary lists its entries in what is called alphabetical order. This is based on the written graph sequence of "a, b, c....x, y, z". But this is definitely not the manner in which morpheme meanings are stored in a speaker's semantic system. In the course of a conversation, say about sports, a speaker desiring to discuss fielding positions in baseball undoubtedly does not proceed through an alphabetical list to come up with center field, first base, left field, right field, second base, shortstop, and third base. Rather, under the domain [baseball], she will search in the sub-domains infield and outfield. Clearly a domain approach appears to be a much more useful model, although the exact nature of semantic cataloging is not yet understood.

¹² Another limitation, which will not concern us here, is the typical dictionary entry's inclusion of word origin, or etymology. This information is clearly not present in a speaker's semantic system unless a morpheme is a very recent loan and therefore part of its significance is its "foreignness."

Describing Domains and Domain Membership

As one might expect, there are a large number of domains for any group of humans. Since these organize and provide the focus for much human interaction and concerns, describing them is as much a task for the cultural anthropologist as it is for the linguist. Among the Nuer, referred to above, the anthropologist could not avoid describing their domain involving “ghok” or cattle. As the anthropologist E. E. Evans-Pritchard wrote, “they are always talking about their beasts. I used to despair that I never discussed anything with young men but livestock and girls, and even the subject of girls led inevitably to that of cattle.” (1940:18-19). In this regard, one of the linguist’s tasks would be to identify dominant domains—those which served as the foci for much speech behavior. However, since this is really still the work of cultural anthropology, it would be more interesting for the linguist to examine the internal structure of domains. One approach is to describe the allocation of different categories of morphemes, and also the co-occurrence of particular morphemes, which co-occur within domains. A domain such as baseball would include nouns (bat, base, pitcher) as well as verbs (slide, steal, retire), adjectives (split-fingered, bloop) and locative (outta there, in the dirt, over the plate). These types could then be examined as to their patterns of use. Some nouns, for example, would be used primarily as objects (“base”), certain nouns would not occur as subjects with certain verbs (*“The shortstop steals second” is not used but “the shortstop makes the throw” is used), certain adjectives occur only with certain nouns (“splitfingered fastball” is used but *“splitfingered triple” is not). In other words, the linguist might work out a “map” of the domain in addition to a listing of its contents.

Sequences

Another approach is to determine if the morpheme meanings of a domain are organized according to a particular pattern. One kind of domain pattern is a **sequence**. Numbers and letters are domains organized in this way. This is not to state that they cannot be utilized in a non-sequential fashion, but that speakers ordinarily use them by sequence. Numbers are usually employed in the sequence “one, two, three....” and letters “a,b,c”. Some other possible sequential domains are ritual (e.g. wedding: courtship—proposal—announcement—invitation list—church service—reception...), performances (e.g. football game: warm-ups—coin toss—kick off...), and certain technical series (e. g. geologic eras: paleozoic—mesozoic—cenozoic or, of more interest to anthropology,

traditional cultural stages of ancient Europe: Mousterian—Perigordian—Aurignacian—Solutrean—Magdalenian).

Taxonomies

Another type of pattern is a **taxonomy**. Unlike sequences, the taxonomic relationship is necessarily one of hierarchical similarity. The meaning of one element in a taxonomy is very similar to other items in certain ways, but different in terms of degree of inclusion—morphemes which are “higher” in the taxonomy also include the meaning of those morphemes which are below them but “lower” morphemes do not necessarily include the meaning of higher morphemes. The meaning of the morpheme *Mammalia* in Figure 4–1 (partial taxonomy of the animal kingdom for primates) includes the meanings of both *Rodentia* (rodents) and *Primates*. However, each of these latter terms really doesn’t contain all of the meaning of the higher term *Mammalia* which would also include *Insectivora* (insectivores). Further, each of the lower terms, since they are in different lines under *Mammalia*, excludes the other (Rodentia does not mean all that is meant by *Primates*), yet they are still similar as members of the same higher inclusive meaning of *Mammalia*. The differences among the elements of a taxonomy are progressive and incremental, and as one enters and proceeds through a taxonomy from more to less inclusive (“higher” to “lower”) the number of elements gets larger and larger. Figure 4–1 shows that part of the taxonomic domain for mammals which includes the human species. Note that this is just one of several formats for depicting a taxonomic domain, with the most inclusive term at the top and the less inclusive terms branching toward the bottom. Starting at the more specific terms—the right in this layout is the bottom, and moving to the left, or top—humans are kinds of hominids (*Hominidae*); humans, gorillas, chimps, and gibbons are kinds of hominoids (*Hominoidae*); the hominoids along with baboons (Old World monkeys) and spider monkeys (New World monkeys) are kinds of anthropoids (*Anthropoidiea*); anthropoids along with tarsiers (prosimians) are kinds of primates; and primates along with insectivores and rodents and many other orders are kinds of mammals¹³.

¹³ This taxonomy is out of date. Current versions would place chimps with humans as Homininae, and place both Pongids and Hominins together as Hominids, among other refinements. Taxonomies are cultural constructs, and change as new viewpoints and data (in this case, fossil and genetic) become available.

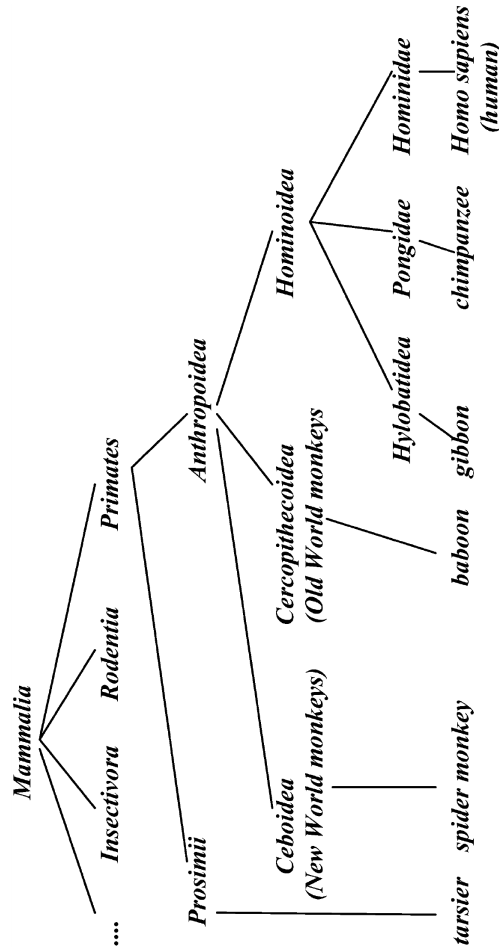


Figure 4–1 Primate Taxonomy
(After Poirier)

Synonymy and Antonymy

Finally, there are many other semantic relations of interest to linguists. **Synonymy** is the relationship of similarity among a group of meanings such that they will be able to substitute for each other in most contexts. There is often a difference of extended meaning among synonyms, in particular one of speech style (for example, formal “please be quiet” and

informal “shut up” would be synonyms but would probably be used in different speech contexts). **Antonymy** is the relationship of opposition. This still involves some similarity on several attributes but opposition on one particular attribute. Thus antonyms “good” and “evil” are both moral qualities but specifically differ as to their social evaluation, “good” is positively and “evil” negatively evaluated. On the other hand, “white” and “short” are not antonyms even though they are different because there is no specific attribute opposition in the midst of similarity.

Linguists would also look at the existence of restricted arrangements of meanings such as dyads (or pairs) (“wife—husband”, “employer—employee”), and triads or sets of three (“Father—Son—Holy Ghost”).

Exercise 4–10 asks you to analyze a domain that is part of your own experience and usage. See how much of the approaches above you can use. Adding to the wonderful complexity of human semantic systems is the common human use of **metaphor**—the use of a morpheme(s) from one domain to apply to a different domain. There will always be a considerable amount of creative license in using metaphors but they must have some degree of semantic suitability in the transfer from one domain to another. For example, while it would be effective to use the sports morpheme sequence “He scores!” for a winning proposal in a business deal, given the similarity of competitive achievement, it might not be effective as a response for a child’s first steps. There is the sense of achievement, but not the competitive advantage. However the study of metaphorical usage really enters into cultural anthropology since metaphor use must be examined based upon community standards.

Challenge of Semantic Systems

Semantic systems are the most challenging aspects of language for the linguist. Their complexity defies any hope for tidy descriptions. Moreover kinds of semantic arrangements discussed above are not mutually exclusive but overlap in many ways. Morpheme meanings (even just considering core meanings) typically occur in many domains (e.g. “baseball” in a sports domain as well as in leisure, childhood and national identity domains). A normal human speaker, given the extraordinary capabilities of the human brain, can make use of the meanings within his or her entire semantic system in multiple and creative ways. Domain linkage, sound, grammatical function, context, hearer’s reactions, personal experience—all have an effect on the speaker’s generation of meaning and the linguist can only map the outlines, and some of the important features, of this open and productive system of communication.

Terms to Know for Chapter Four

significance (of speech signals)	dictionary model
arbitrariness (of speech signal significance)	denotative-connotative meaning
openness and productivity (of language-based communication)	core and extensional meaning
morpheme	analogy
identification of morphemes	homophones
co-occurrence (of phoneme sequence and gloss)	semantic categories
gloss	semantic relationships
embedded meanings	domain
word	co-occurrence pattern
componential analysis	sequential pattern
	taxonomic pattern
	synonymy
	antonymy
	metaphor

CHAPTER FIVE

GRAMMAR: THE ORGANIZATION OF WORDS AND SENTENCES

In this chapter we will examine how words and sentences are put together. In the last chapter we introduced the concept of a morpheme as a unit of meaning, in this chapter we will look at it as a unit of arrangement (without ignoring its semantic attributes). We will begin with the concept of a word as a minimal spoken unit, and the fact that they are usually made up of combinations of morphemes. Morphemes can be categorized as to the way they are used to build words—as stems and affixes. An interesting result of morpheme combinations within words is the phonemic modifications which morphemes may undergo in certain phoneme and morpheme environments. Thus we shall examine the concept of allomorphs (much as we did for the allophone concept in Chapter 3). In the last part of this chapter we will look at the way in which words are put together to form phrases and sentences. A sentence is not just a sequence of words, but rather a sequence of positions, or “slots,” into which certain kinds of words can appear. We will examine the way these word classes operate and are defined. It will be necessary to go back to morphology at this point since word classes are bound up with certain kinds of affixes, both derivational and inflectional. Finally, we will examine the significance of the focus on syntax and the attempt to write grammars that can generate all of the sentences in a language.

The Importance of Grammar

In the last chapter we examined the semantic system. There we dealt with units having significance to a community of speakers. But this significance, or meaning, can only be expressed with an observable signal that, for our purposes here, is a speech signal. The communication of meaning, therefore, takes place by means of the production and organization of meaningful units of speech. We have seen that the smallest of these meaningful units are morphemes. These either occur by themselves, or are combined with others, to form words. Words are the smallest units which can constitute an utterance (or, to state it in another way, a word is the minimal spoken unit). Phrases and sentences, then, are

made up of words. And discourse, then, is made from phrases and sentences.

However, while phonology may be analyzed to a great degree without considering the significance of the speech signals that its phonemic units—phonemes—form, grammar cannot be studied without taking significance into equal consideration. We have seen that the identification of morphemes depends upon associating them with meaning. Similarly, the analysis of grammatical processes of word and phrase formation depend upon identifying and associating these with semantic processes. In other words, grammar can not be just an abstract calculus of morpheme arrangements. It must at the same time deal with the semantic functions of morpheme arrangements.

Morphology

Morphemes and Word Structure

The first grammatical process we shall examine is **morphology**, or the system by which morphemes are combined to form **words**. As we have seen, a morpheme is the minimal unit of significance. It is made up of one or more phonemes and it possesses some kind of significance such that its presence or absence alters, in some way, the total significance of a phrase or sentence. It cannot be divided (completely, with nothing left over) into smaller phonemic units with significance. For example, the English word “farmers” /fájmejz/ can be broken down into the morphemes /fájm/, /ej/, and /z/ (we will take this as given here and not supply a morphological analysis). Therefore /fájmejz/ is not a minimal unit and is not, by itself, a morpheme. /fájm/ on the other hand, is a morpheme since it cannot be divided. The reader, however, might notice that the beginning segment of /fájm/, */fáj/, is familiar and perhaps is the morpheme “far” meaning [a great distance]. However, this division is not valid since it leaves the remainder of /fájm/, /m/, unaccounted for. /m/ has no significance within the sequence /fáj/, in particular it has no significance that is related to the /fáj/ sequence with the meaning [a great distance]. Thus, while there is a morpheme which consists of a /fáj/ sequence, it is not present in /fájm/. The initial three phonemes of /fájm/ *by coincidence* are the same as what would, *in another context* (e.g. “Is it far?”), be a morpheme. But in the /#_m / context (or frame) they are not a morpheme but are inseparable from the /m/ as the four phonemes of the morpheme /fájm/.

Stems and Affixes

Semantic Distinction

Morphological analysis is based upon classifying the morphemes which constitute words into two types—**stems** and **affixes**. This classification is necessary since the stem is the base, or the morphological landmark, to which the affixes are added in a particular manner, singly or in combination. For example, in each of the following forms /kæt/ “cat” is the stem and all other morphemes are affixes:

- | | | | |
|-----|-----------------------------------|-------------|---|
| 5.1 | /kæts/ | “cats” | [more than one cat] |
| 5.2 | /kæti ¹ / | “catty” | [having the characteristics of a cat, in particular as applied to humans being malicious] |
| 5.3 | /ən <u>kæt</u> la ¹ k/ | “uncatlike” | [not having the actual characteristic of a cat] |

How is this classification done? It is not as easy as it appears, and is based on a number of criteria. One criterion would be the greater semantic “weight” (or contribution) of stems to the resultant meaning of the word. In forms 5.1 to 5.3, /kæt/ appears to contribute more to the total meaning of the resultant words than /s/, /i¹/, /ən/, or /la¹k/. These latter morphemes certainly contribute to word meaning, but more as adjustments to the morpheme /kæt/ than as significant determinants of word meaning. In other words, they each just seem to modify [cat] in various ways. However, in item 5.2 a case could be made that /i¹/ has a semantic contribution as great as /kæt/ since the resultant meaning is quite different from a simple [cat] plus [(modifier)]. (Another way of looking at this is to claim that /i¹/ has the primary significance of [quality of] and so is the stem while /kæt/ only indicates which quality is applied.) In other languages, too, this criterion of relative semantic weight may not be easy to apply. For example, the Eskimo word for “I am quickly making you a little box for percussion caps,” /qanuɟayruŋəstaɟutaxaŋlaraxkiyakamkən/, or, lined up with each morpheme’s approximate meaning:

/qanuɟay –ruŋ –əsta –ɟu –tax –aŋ –la –raxki
 (brass –thing –small –collective –container –small –make –quickly –
 –ɟa –ka –m –kən /¹
present –indicative –I –you),

¹ The dashes have been added to assist the reader’s understanding of the morphological structure of the word. They have, of course, no phonemic value. This example is from Atkinson, et al., p. 138.

challenges a purely semantic decision on which is a stem and which is an affix. Is the stem /run/ “thing” or /tax/ “container,” to use those morphemes which seem from the gloss to be substantive (i.e. what we might be disposed to think of as “nouns”)? Or is it /la/, which appears to be a “verb”? Or as another example, consider the much shorter Kiswahili (East Africa) word /ki-le/ “that” (as in /kile kiti/ “that chair (which I just mentioned)”). The latter morpheme, /le/, carries the meaning of [referred to] and would seem to therefore be the stem, but the former morpheme, /ki/, actually indicates that a particular following word, e.g. /kiti/ “chair,” will be specified in some way, in this case as something already mentioned, and so it may lay claim to being a stem semantically.

Other Distinctions between Stems and Affixes

In fact, the semantic weight criterion is just one of five diagnostic clues to the identification of stems. Another is the **frequency test**. In any language, there will be relatively few affix morphemes and a very large number of stem morphemes. Imagine that affixes are the “nut and bolt” morphemes of word formation—morphemes which fit stems in various ways to function as words in phrases. There are relatively few types of “nuts” and “bolts” needed to do things such as pluralization, case, or tense. but these are used over and over again with all the different reference morphemes of entities and activity. In regard to identifying the stem in the Kiswahili word /kile/, referred to above, /ki/² occurs in hundreds of words with other morphemes (e.g. /kímo/ “inside,” /kíko/ “thereabouts,” /kíti/ “chair,” /kízúri/ “good”), while /le/ only occurs in one word type. Thus /ki/ is more usefully considered as the affix, while /le/ is treated as the stem. Therefore in a large corpus (one that has items from a variety of domains) those morphemes which occur infrequently will probably be stems and those which occur frequently should be affixes.

A third difference is that while the semantic relationship (see Chapter 3) among stems can be quite diffuse (i.e. there need be no special connection among them), the relationship among affixes is usually one of synonymy or antonymy. For example, for the English “wings,” “nails,” “posts,” “oxen,” and “alumni” there is no special meaning relationship among the morphemes “wing,” “nail-,” “post,” “ox,” and “alumn,” but there is the semantic relationship of synonymy among “s,” “en,” and “i”—that of [plurality]. Going back to the “nuts and bolts” idea, the affixes of a language perform a relatively small number of modifications on a very large number of different stems whose semantic relationship may be

² The high stress turns out to be unnecessary for the morpheme’s form, as the following examples show.

nonexistent (not always—some stems may be synonyms). So if there are morphemes in a corpus which seem closely related along some semantic continuum these may likely be affixes, and those which have little or no relationship are probably stems.

A fourth difference is that stems *tend* to be longer than affixes in regard to number of phonemes. And a fifth is that stems *tend* to be free morphemes—those which can stand by themselves as words, while affixes will be bound morphemes—those which only occur with other morphemes in words. These latter two criteria are not strong distinctions and would never be sufficient all by themselves to identify stems. But they can assist the linguist in starting on the problem of identifying the stems and affixes of another language.

Therefore, after identifying the morphemes in the items of a corpus, the linguist would start making decisions about how to allocate these into a stem group and an affix group. For example, she might decide that certain of the morphemes appear to contribute more weight to the item's total meaning (based upon the gloss, but better still would be native speakers' judgements) than any of the other morphemes in that item. This morpheme would tentatively be placed in the stem group. This test would be applied to all morphemes in all items. In addition, she would apply a frequency test and a semantic relationship test to confirm this tentative placement. With luck the other tests (and perhaps the length and free/bound test as well) would confirm the initial placement that certain morphemes in the corpus are usefully considered as stems and the rest would then be affixes. It is never as easy as this, often different tests contradict one another, but this is the procedure. At this point you should try to apply what you have learned by working out Exercise 5-1 and figure out from the Klingon corpus (used in Exercise 4-7) which morphemes are stems and which affixes.

Kinds of Affixes

Affixes are classified by where they are joined to a stem. **Prefixes** join with a stem at its beginning, **suffixes** join at the end, and **infixes** are inserted into the stem at some definite internal position after the initial phoneme and before the final phoneme. English has several prefixes (“unhappy”, “retry”, “asocial”) and many suffixes (“cats”, “farmer”, “national”, “walked”). But for infixes we must look at another language. For example Bontoc (Philippines) has an infix /-um-/ as in the word /fumikas/ “he is becoming strong” (compare /fikas/ “strong”). Or Arabic with its many vowel infixes, such as occur in /zoker/ “one who remembers”, in which the stem is /z-k-r/ “remember” and different vowels are inserted as infixes, in this case /-o-e/ “one who.”

Words usually possess only one stem³, but commonly can have more than one affix. The Kiswahili word /kičéko/ “joke,” for example, has both a prefix, /ki/, and a suffix, /o/. The meaning of these are not easy to briefly define, but /ki/, the same prefix we saw in /kile/ above, is a noun case prefix which also means [singular] (compare /vičéko/ “jokes”) and /o/ is a noun-forming suffix (compare /čeka/ “(to) laugh” having the verb-forming suffix /a/. Or the Turkish word /kollarinizdan/ “from your (plural) arms” which is constructed from a stem (/kol/ “arm”) and four suffixes (/lar/ [plurality applied to the stem meaning], /in/ [second person pronominal reference], /iz/ [plurality applied to the pronoun], and /dan/ “from”). There are also **obligatory affixes** as well as **optional affixes**. Obligatory affixes are those which must occur as a part of the formation of certain words while optional affixes may or may not occur. English, for example, has few obligatory affixes but mostly optional affixes—there are few word types in English in which certain affixes must occur. In Kiswahili, on the other hand, there are a number of word types which have obligatory affixes. For example, the Kiswahili word represented by /kile/ “that” (see above page 158) consists of the /le/ stem and one of many obligatory prefixes: /ki/, /vi/, /yu/, /wa/, /zi/, In other words, this kind of word cannot occur without an affix (and in this case both the affix and the stem would be bound morphemes).

Affix Orders

Morphology, therefore, involves the categorization of morphemes into stems and affixes, affixes into prefixes, infixes and suffixes, and prefixes and suffixes into **affix orders**. This last step must be done when more than one prefix or suffix can be attached to a stem (infixes apparently only occur singly). An affix order is a space, or slot, in a sequence of positions, starting from the stem, which an affix has the privilege of filling when it is present. For example, in the Turkish word given above, /kollarinizdan/, /lar/ is a first order suffix because when it is present it has the privilege of occupying the first suffix position following the stem, /in/ is a second order suffix because it has the privilege of occupying the second suffix position following the stem when a first order suffix is present (if no first order suffix is present then it is next to the stem). An order usually contains more than one affix, and obviously only one of these can occur at a time in this particular position. Furthermore, the affixes belonging to each order share the same general semantic reference. The fourth suffix order for Turkish nouns, for example, may contain either /dan/ “from,”

³ There are exceptions to this. See stem compounding, below, as a means of word formation.

/da/ “in,” /a/ “to,” or /li/ “with” all of which are locational morphemes. Similarly the second order suffixes all refer to person (e.g. first person, second person). Of course, affix orders are only determined after a sizeable number of words are studied and analyzed. This is because representatives of all affix orders typically do not occur in every word, and some may occur in only a few words. It is not, therefore, important to identify the particular order number of an affix so much as it is to establish its place in a sequence. In Turkish nouns, for example, a small number of nouns should establish that the person number suffix follows the noun number suffix and precedes the person number suffix. And this is more important than attaching the numbers 1, 2 and 3 to these suffix orders. Additional words in the corpus may yield different suffixes which would change the numbers, but not their relative sequence. Exercise 5–2 will allow you to figure out affix orders for Turkish nouns.

The Concept of a Zero Morpheme

There is one interesting aspect of affix orders which has created some controversy, this is the concept of a **zero morpheme**. Imagine a series of four prefix orders for the verbs of some language with the second order referring to tense. There are second order prefix morphemes for [past] and [future], but none for [present]. The linguist determines that when there is no second order prefix, just members of the first, third and fourth orders, the appropriate reference is to a present tense (not merely the absence of a tense). Consequently, the linguist concludes that the absence of a second order prefix (i.e. a “zero” morpheme) refers to the present tense. This would seem to be a useful conclusion, especially for languages with a morphology requiring the presence of several affix orders on certain classes of words such that the **absence** of any affix for one of the orders can easily serve as a signal—a zero morpheme. But for languages with little required affixation, where the absence of an affix is not a clearly marked signal the zero morpheme concept may not be useful. English nouns, for example, have only one inflectional⁴ suffix order containing one suffix morpheme meaning [plurality] (/s/ in many occurrences, but see allomorphic variation below), for example “cat + s” meaning [more than one cat]. It seems awkward to view the absence of this morpheme (“cat _”) as a zero morpheme meaning [singular]. On the other hand, if we do not use a zero morpheme for [singular] then we are left with incorporating [singularity] into the “cat” stem allowing us to use “cat” by itself and also impart the meaning of [singular]. If we assume that there is a suffix for

⁴ The term “inflection” will be discussed later in the chapter. There are other noun suffixes but these are classed as “derivational”, also discussed below.

[plural] but no suffix for [singular], not even a zero suffix, then given that the [singular] meaning is nonetheless an active part of how the noun form is used⁵, where is it to be located morphologically? The answer would appear to be to place it in the noun stem. Thus the morpheme “cat” not only has the meaning of [a kind of animal] but also [singular]. Multiple meanings for one morpheme are certainly common, but what, then, do we do with the [singular] part of the stem’s meaning when the [plural] suffix is attached. It would seem that we must then have some rule with the effect that when the [plural] suffix is present the [singular] aspect of stem meaning disappears. Is the concept of a zero morpheme much more complex than this? If nothing else, this illustrates the intricacies of language systems, and the challenges which await linguistic analysis. [Note: that aspect of word formation dealing with compound stems is discussed following the section on inflection and derivation]

Variation in Morpheme Shape

Allomorphs

At this point we can undertake a reexamination of the process of morpheme identification. We saw in the preceding chapter that morphemes can be identified in a corpus by determining the co-occurrence of phonemes and gloss (i.e. meaning). This association between phonemic form and meaning (as approximated by the gloss) was presented as being absolute. This, however, is not true. While meaning must remain constant, the phonemic form of morphemes is often variable. An example of this phonemic alteration can be seen in those English morphemes for plurality which, though spelled “-s”, occurs in speech as /s/, or /z/. (The plural morpheme, in fact, has other forms than these and they will be discussed below.)

Thus “cat, cats” /kæt, kæt̚s/ and “dog, dogs” /dag , dagz/. A linguist studying English and having these forms in a corpus would be puzzled by the recurring gloss of plurality for the second of each pair of words in the absence of a common recurring phonemic shape. For some glosses of plurality the shape is /s/ yet for others it is /z/. One possible explanation is that they are, in fact, different morphemes but have meanings which stand in the relationship of synonymy. The fact that the gloss is the same, as was discussed in chapter 4, may simply indicate an incomplete translation or

⁵ In other words, we cannot say that [singular] is not present since it has an affect on other words. Compare “the cat_run home” and “the cat̚s_run home.” Similarly, certain articles and demonstratives only occur with singular noun forms. For example, compare “a cat” and “the cat̚s”.

investigation of each morpheme's complete significance. Two kinds of data will make this explanation untenable. First, the informants will insist that the meanings are the same, i.e. that "cats" and "dogs", whatever other differences in meaning are present due to the presence of "cat" and "dog", respectively, each has the significance of plurality. It is *not* the case, for example, that /s/ would mean [two or more definite things] while /z/ would mean [two or more indefinite things]. The second is the finding, following the collection of many items and carefully examining the environments in which each shows up, that these occur in a complementary distribution. The existence of complementary distribution, along with sameness of meaning, is a sure indication that these are **allomorphs**⁶, or different phonemic shapes of the same morpheme.

Phonologically (Phonemically) Distributed Allomorphs

This type of allomorphic variation is fairly straightforward—often referred to as “regular”—and native speakers (speaking their first language) are typically not consciously aware of these alternations of morpheme form. In these cases, a rule can be stated which accounts for the variation in morpheme shape according to certain characteristics of its phonemic environment. In other words certain allomorphs of the morpheme occur in one particular set of phoneme environments and other allomorphs occur in other environments, thus the allomorphs are in phonemic complementary distribution. (This is the same situation as for the variation in allophonic variation according to phonetic environments described in Chapter 3.) In our example for English plurals, let us assume that the underlying or base form for the plural morpheme is an alveolar fricative. It is a suffix and so will be placed next to the final phoneme of the stem. If the final stem phoneme is a voiceless stop, then the phonemic shape of the plural suffix will be realized as a voiceless alveolar fricative. If the final stem phoneme is a voiced stop then the plural suffix will be realized as a voiced alveolar fricative. Thus the plural morpheme suffix occurs as /s/ when attached to “cat” /kæt/ (becoming /kæts/) and as /z/ when attached to “dog” /dag/ (becoming /dagz/). As you might be realizing, the rule for English plurals needs to be a little more comprehensive and Exercise 5–3 on English plural allomorphs will allow you to formulate a more complete rule for the English plurals that are in

⁶ This term should be somewhat familiar to the reader due to its similarity to the term allophone. Both terms refer to variants of a base form, allophones are variants of a phoneme and allomorphs are variants of a morpheme. The allo- prefix is used in several words to indicate such variations.

phonological distribution, including a third plural form /əz/ as in “roses” /jɒ^wzəz/.

This process, where variation in morpheme shape is in accord with phonological conditions, has often been referred to as **morphophonemic adjustment**. This is viewed as the result of three conditions: The first is the presence of a base form of a morpheme which has specified (if generalized) phonemic features (as was done above in stating that the base form for the English plural was an alveolar fricative); second a system of phonemic restrictions on phoneme sequences within a syllable (such as the formula for English syllables in chapter 3); and third the morphological processes of binding morphemes together within words, thereby placing the phonemes at their boundaries together and adding to existing syllables. One implication of a morphophonemic account is that the phonological system dominates the morphological system in that its guidelines for appropriate phoneme combinations take precedence over the morpheme shape. Consequently when two morphemes are placed together creating an inappropriate combination, the phonemic shape of one of the morphemes (or both) is altered so that the phonemic rules are not violated.

Morphologically Distributed Allomorphs

There is another type of allomorphic variation, one which is not related to phonemic environment. Still based in complementary distribution, but in this case the environment is morphemes rather than phonemes. In other words, one allomorph occurs with certain morphemes and another allomorph occurs with other morphemes (none of these allomorphs, by the way, violates phonemic restrictions). For example, the variation among the other allomorphs of the English plural morpheme, for example /ɪn/ (“oxen”), /ə/ (“data,” singular “datum”), /ə/ (“syllabi,” singular “syllabus”), cannot be accounted for by phonological rules. Instead they must be associated with particular morphemes (in this case stems): for example /ɪn/ occurs with “ox”, /jən/ with “child”, /ə/ occurs with “dat” and “phenomen,” and /ə/ occurs with “syllab” and “alumn”. The reader can check the phonemic shape of these stems to confirm that if they operated with the phonologically distributed series, they should occur with the /s, z, əz/ allomorphs (i.e. “ox” /aks/ would have been */aksəz /, “syllab” would have been */siləbz/). On the contrary these particular stem morphemes “possess” their own sets of plural suffixes. (Of course, in addition, many have singular suffixes as well, such as “datum”, or different stem allomorphs, such as “child” /tʃaɪld/ for the singular but for the plural “children” this is /tʃɪldjən/.)

Morphologically distributed allomorphs are often referred to as “irregular” because their occurrence cannot be reduced to a set of phonemic conditions (i.e. a “rule”). Instead, each of these allomorphs must be listed with its controlling morpheme. Interestingly, the presence of these kinds of allomorphs has sometimes given rise to odd beliefs about the character of a language. The use of the descriptive label “strong” for those English verbs whose allomorphs in different grammatical cases (i.e. tense) was marked by changing the verb stem vowel rather than by adding a suffix (which was labeled “weak”), “dive—dove” and “run—ran” rather than “type—typed” or “drop—dropped.” A colleague once informed the author that German was a “strong” language because of the presence of these kinds of forms (I think that the implication was that it took more dedication to learn). In fact both kinds of allomorphs are “regular” insofar as they have specified distributions, they do not occur randomly. Further the “strong-weak” labels have no bearing on the nature of the speakers and their language-acquisition. It does appear that phonologically distributed allomorphs may be acquired earlier by speakers since the phonemic system is mastered by age 5 or so. Morphologically distributed allomorphs are acquired, as a group, later simply since they may often be less frequently employed. Consequently English children are likely to use “oxes” instead of “oxen” or “alumnuses” instead of “alumni.” All they are doing is applying the phonological rather than the morphological rule, which they will acquire as they learn it. In some cases acquiring the morphological rule is attached to a special experience, in the U.S. formal education, and so failure to use the morphological rule when it is appropriate is associated with a pejorative assessment of the speaker—being childish or uneducated. But there is no support for the conclusion that one kind of allomorphic rule is “better” than another. This was in part the issue with the varying forms of English such as Black English Varieties that employed forms as “he **go** home yesterday” instead of “went” and were then injudiciously labeled as lacking in mental capability. It is simply that these morphologically associated forms are taken too seriously by some speakers. However, all languages possess some allomorphic variation of both kinds among their morphemes. To this extent, this merely indicates that languages are dynamic and malleable systems.

Allomorphs and Language History

The question should arise as to why this variation in morpheme shape occurs. No living language is a perfectly balanced, harmonious closed system. Given the nature of human life, a language must be a dynamic, flexible system which is continuously subjected to change (which is, to a

greater or lesser degree, both resisted and desired). Our human language facility apparently requires that we produce sound signals that are guided by both sound and grammatical designs. In other words, we do not operate randomly. However, over time (subsequent generations of speakers do not precisely reproduce the speech patterns of the preceding generation) and with changing social circumstances (e.g. migration, contact with other groups, population changes) phonological and morphological changes accumulate. Change never occurs quickly or universally throughout a language, nor does change come for the same cause. The current state of English plurals is the result of over a thousand years of various kinds of changes. One was the change from a system of different noun cases for which there would be a number of suffixes for plural depending on how the noun was being used (for example whether being used as the subject of a sentence or as the object) to a system where these case endings disappeared from usage and one was used for all plurals. Another was the various conquests of Britain by Germanic speakers (Angles and Saxons) and by French speakers (the Norman Conquest) who languages were the source of new words and plural forms. The sub-discipline of Historical Linguistics, referred to in chapter 1, deals with the challenges of reconstructing these kinds of changes to gain a long perspective of the history of a language system.

Syntax

The Semantic Structure of Sentences

A word, consisting of a single morpheme or some combination of stem and affix(es), may constitute a complete utterance; but it will more likely be combined with other words. **Syntax** is the part of grammar which organizes words into phrases and sentences. An important organizational principle of syntax is the grouping of words into different classes according to how they contribute to the organization of the sentence. One attribute of this contribution is semantic significance. As we saw in the previous chapter, one part of a semantic analysis is the classification of morphemes (or words) according to their meaning. Nouns, for example, had meanings representing some kind of entity, verbs, some kind of activity, and so on. Let us assume that this kind of semantic classification is possible and that all words in a language are grouped as nouns, verbs, adjectives, or adverbs. These are traditionally called parts of speech and provide one way of classifying words by their syntactic function. A **sentence** must be examined not just as a sequence of words, *but as a particular sequence of parts of speech*. The last phrase in the preceding sentence can be described as a sequence of: conjunction, adverb, article,

adjective, noun, preposition, noun, preposition, and noun. This approach permits certain kinds of descriptive statements concerning how parts of speech may be combined. For example in English, adjectives precede nouns, and articles (or demonstratives such as “this”) precede adjectives. Put another way, within an English noun phrase the semantic aspect of [specification] contributed by articles and demonstratives is produced first by a speaker anticipating a following noun [entity]; the semantic aspect of [quality, attribute] contributed by adjectives is produced by a speaker after an article [specification] but still prior to a following noun [entity]; and the semantic aspect of [entity] contributed by nouns is produced last. Other languages might have different arrangements of parts of speech. For example, Kiswahili has a standard sequence of noun, adjective, demonstrative, the opposite of the English sequence. However, as interesting as the use of semantic attributes is for analyzing syntax, it has two drawbacks. The most important of these is that semantic contributions of words do not completely define sentences. The other is that speakers are able to understand and use words whose meaning is obscure or unknown, and therefore their part of speech membership could not be known.

Defining a Sentence by Prosodic Features

The first problem concerns the definition of sentences. A sentence is difficult to define using the semantic attributes of its constituent words by themselves. Instead, a **sentence** can be more usefully defined as an intonation contour. We saw in chapter 2 that one of the important features of speech was stress and pitch (or relative loudness and relative tone). Whatever segmental phonemes must be strung together in order to complete the words which will constitute a complete utterance, a certain sequence of stress and/or tone levels must also be produced if the word sequence is to be accepted as a unit of discourse. For example, following Charles Hockett, a typical English statement consists of the following intonation sequence / ²³¹ ↓ /.

Each raised number represents a combined stress and pitch level (in English stress and pitch are inseparable) and the vertical arrow represents either a slight rise or fall of pitch, in this case a slight fall. Thus “John go get the cake” / ²jan go^w get ðə ³keɪk¹ ↓ / qualifies as a complete utterance as one part of a discourse. Mariana and Louise could be discussing party arrangements within earshot of John:

Mariana “I guess that takes care of everything.”
 / ²a¹ ges ðæt teɪks keɪ əf ³evjɪθɪŋ¹ ↓ /

- Louise “Yes, except for the cake.”
 /²yes eksept fə^wj ðə³ke^lk¹ ↓ /
- Mariana “Right, I almost forgot.”
 /²ja^t a¹ almo^wst fə^wj³gat¹ ↓ /
 “John, go get the cake.”
 /²jan go^w get ðə³ke^lk¹ ↓ /
- Louise “We just have time to mix the punch”
 /²wiⁱ jəst hæv taⁱm tu^w miks ðə³pənč¹ ↓ /

Each of these utterances consists of a /²³¹ ↓ / intonation contour (which the reader should check by saying them as naturally as possible). Moreover, according to Hockett, it is the particular sequence of / —³¹ ↓ / which usually signals the end of a complete utterance (there are others). The two speakers, Mariana and Louise, in fact should take turns speaking on the occurrence of this sequence, otherwise they would be “interrupting,” as illustrated in the following exchange:

- Mariana: /²jan go^w get /
- Louise: /²wiⁱ jəst hæv taⁱm tu^w miks ðə³pənč¹ ↓ /

Louise has started talking prior to an occurrence of a / —³¹ / by Mariana forcing her (in this case) to stop, undoubtedly feeling frustrated at not finishing her utterance.

The implication of the use of an intonation frame to define a sentence, or, better, a completed utterance, is that the hearer is cued to treat something with a /²³¹ ↓ / intonation pattern as a complete utterance even if its apparent morpheme content is not understood. For example, /²mal grum budrik oⁱ klin³ farb¹ ↓ / will at least be considered as a possible sentence even though (to the writer’s knowledge) it consists entirely of nonsense syllables, since it has the appropriate intonation of a sentence such as “the boy broke that small toy.” Still without any familiar morphological landmarks the hearer will be confused by the strangeness of the apparent morphemes and the familiarity of the intonation. A more likely example would be “²Kim tweeted that app mucho ³pronto¹ ↓” This will confuse the unfortunate hearer who is unfamiliar with the web-based cell phone application of simultaneously sending a message to a number of people (“tweet”), the abbreviated term for a kind of application that can be performed by an electronic device (“app”), and the quasi-Spanish jargon for quickly (“mucho pronto”). However, in this case, the hearer is undoubtedly convinced that she has heard a complete utterance because the appropriate intonation is accompanied by three familiar morphemes

(“Kim”, “ed”, and “that”). This is why Lewis Carroll’s poem “Jabberwocky” (from *Through the Looking Glass*) can be “understood” or at least appreciated by a listener—it contains enough familiar morphemes so that, coupled with an appropriate intonation, it seems as if it should be understandable:

²**Twas** ³brillig¹ ²**and the** slithy toves
 did gyre **and** gimble in the ³wabe¹
²**all** mimsy **were the** borogroves
and the mome raths out³grabe¹... (boldface added)

The structure of sentences in each language, therefore, must be examined within these intonational frames which mark complete utterances.

This last point leads into the second problem in a purely semantic approach to sentence structure. It is revealed in the ability of speakers to use words or morphemes in acceptable utterances without knowing what they mean. Using the above examples, a speaker could hear the unfamiliar words and still be able to respond with acceptable utterances such as: “I’ve never tweeted anything”, “which app was that?”, or “Are toves always slithy?” and “How big are borogroves, anyway?” It is as much the manner in which words (such as “tweet. app, toves, slithy and borogroves”) are used in combination with familiar words (such as “that, the, and were”) which determines their places in a sentence as it is their meanings. This is not to suggest that meaning is unimportant, it undeniably is. It is just not sufficient, by itself, to explain syntax. A sentence is a sequence of words providing structural information as well as semantic reference. We shall come back to this topic below.

At this point we will return to the question of how to investigate the organization of sentences by establishing “parts of speech” as categories of word types, or **word classes** (the more common term). First, of course, we will study word groups having appropriate intonation. These can be determined by careful observation of the intonation of word sequences judged to be complete utterances by speakers. A corpus consisting of sentences (not just words—although a sentence may consist of a single word) will then be obtained, and this should contain a sufficient number of common references so that the same words (or morphemes) will occur in many corpus items.⁷

⁷ One way of achieving this is to encourage the linguistic informant to discuss a series of related activities, such as butchering and cooking a chicken, cultivating a field, or taking care of a child.

Word Classes

Syntax is that part of grammar which deals with the organization of sentences. We have already defined the sentence as an intonational unit. While morphemes are the ultimate components of sentences, in the sense that the semantic function of a sentence is dependent upon the morphemes it contains, the effective components of sentences are words. Therefore, describing the organization of sentences is essentially describing how words are organized together within sentences.

The first step in analyzing sentences is to group words into **word classes**. A word class consists of all those words which can substitute for each other in the same position in a sentence. For example, in the English sentence "The young boy ran home." there are many words which can substitute for "boy": "girl, man, dog, . . ." There must be some semantic appropriateness, of course. Not all members of a word class can substitute in every context; for example, "submarine" will not sensibly substitute for "boy" in the above sentence, but it can in the sentence "Jane looked at the boy." Whereas the words "of, the, rapidly, cry, pretty..." cannot substitute for "boy" or "submarine" in any sentence (i.e. *"Tom looked at the rapidly.>"). Consequently "boy" and "submarine" can be placed together in the same general word class, but not together with "rapidly."

Word classes can be defined in three ways. The first, but not necessarily the most useful, way is by semantic reference. "Boy" and "submarine" are commonly considered to be members of the English word class "noun" because they are persons or things—i.e. the elementary school dictum, "a noun is the name of a person, place, or thing." We saw in the previous chapter that one part of semantic analysis is to group morphemes by the nature of their semantic reference and this is one important area in which semantics and syntax overlap (although here we are classifying words and not only morphemes). Unfortunately, given the complex imprecision in the open and productive nature of human meanings, semantic definitions are often insufficient by themselves to clearly distinguish one word class from another (e.g. in English "running" as an action belonging to a verb word class and "running" as an action belonging to a noun word class or even as an adjective word class).

The second way of defining a word class is by its distinctive pattern of occurrence in regard to other word classes. For example, "boy" belongs to a word class which can be preceded by a word like "young" and followed by a word like "ran." In other words, only members of a certain word class can occur in the frame "The young ____ ran home." An English "noun," therefore, is a word which can occur preceded by a member of an adjective word class and followed by a member of a verb word class. Even though this is a circular definition (an "adjective," on the other hand, is a word

class which can occur followed by a “noun”), it does state exactly one of the important aspects of syntax—word class order. Just as we had to take into consideration the order of affixes in word construction, we must in turn study how word classes are ordered in particular sequences in sentences. Sometimes referred to as **syntagmatic** structure, the pattern of how language units are put together in sequence (from phonemes to word classes and phrases) is an intrinsic aspect of the way humans use language. Thus in the sentence “That big house was sold.” we see a common English noun phrase order: article + adjective + noun + verb phrase. In Kiswahili, however, the sequence would be noun + adjective + article + verb phrase (/ nyumba kubwa ile iliuzwa/ “house big that (it) was sold”). These sequence rules can be complex, for example in English if more than one adjective is used these must be furthered ordered in a certain sequence according to semantic type, “big old green house” (i.e. [size] + [age] + [color], but not some order such as **“green big old house”*). By the way, this last example demonstrates that while semantic content may not always be the most useful criterion for defining syntactic units, it certainly cannot be ignored, and is often indispensable.

Finally, word classes can be defined by the presence of special morphological structure. This may include distinctive types of affixes (e.g. word class X has prefixes while word class Y has suffixes) as well as, and much more likely, distinctive affix sets. In Swahili, for example, nouns take a certain group of singular and plural prefixes (/m-, wa-, ki-, vi-, mi-, n-, ji-, ma-, u-, pa-, ku-, 0-(zero)/ while verbs take a different group of object, tense, and subject prefixes (tenses: / na-, ta-, me-, li-, si-, ka-, ki-, hu-, nge- / among others). The careful reader will notice that there is one homonym⁸ (/ki-/ and there are others as well not included) demonstrating that it is the entire set of prefixes which define different Kiswahili word classes, not individual prefixes. A similar situation occurs in English where, while both nouns and verbs share a suffix homonym, “-s”, the entire set of noun suffixes is different from the set of verb suffixes. The most useful approach would be to use all three kinds of definitions. An English noun, consequently, would be: 1) semantically the name of an entity, event, location or concept, which 2) syntagmatically occurs preceded by adjectives and followed by verbs (in addition to several other frames), and 3) morphologically only occurs with the inflectional⁹ plural suffix {-Z₁}. (This is the designation for the plural suffix.)

⁸ Since these illustrations are for written forms, we will be discussing “homonyms” rather than “homophones”.

⁹ Discussion of “inflectional” affixes is given in the next section.

Inflection and Derivation

Those affixes which occur as part of the words which make up a word class are called **inflectional affixes**. These affixes modify the stems of the word class but do not change their syntactic function (i.e. change their word class). This semantic modification is an important part of the function of inflectional affixes. Common inflectional semantic modifications for verbs are: “tense”—an association of action with time; “aspect”—the duration and completion of actions; and “mood”—who performs actions, under what conditions and upon whom are the actions done. Similarly, nouns are commonly inflected for “number” and “gender”, while adjectives may be inflected for degree of comparison.

Derivational affixes are those affixes which make the resultant word a member of a different word class from the one occupied by the stem. It “derives” a new word. English has many derivational affixes, such as /ej/ “er” which moves a stem from the verb word class (e.g. “run”) into a noun word class (e.g. “runner”) (refer to Appendix F on English affixes). Once a derivational affix has been attached to a stem, it now takes on the inflectional affixes of its new word class (“runner” can now be pluralized with the {Z₁} suffix as “runners”).

Roots

A **root** is like a stem, but it has no original word class membership. Stems, by themselves, without any derivational transformation, can take the inflectional affixes of their respective word class. But roots have no inflectional affixes because they have no word class identity—they must take a derivational affix first, thereby becoming a member of a particular word class, and only then being able to be inflected. For example, Kiswahili has an extensive system of roots: /pend/ is neither verb nor noun but once it receives the derivational affix /a/ it becomes a verb stem /penda/ and can be inflected by tense, subject and object prefixes, for example /nitakupenda/ “I will love you.” However, if the derivational suffix /o/ is added then it becomes a noun stem /pendo/ and can be inflected as a noun for number (singular) /kipendo/ “a loving act” or (plural) /vipendo/ “loving acts”.

Compound stems

Another way of creating stems is to add two (or more) morphemes together which in other constructions would serve as stems by themselves, for example “barnyard.” “Barn” and “yard” are each stems in English and function to make words by inflection and derivation: “barn + s”, “yard +

like”, “un + yard + like”, “barn + ish”. Combined, or “compounded,” however, they operate as if they were simply one unitary stem that can be combined with the noun inflectional and derivational affixes — “barnyard + s”, “un + barnyard + like”, and so on. Note that it would not be likely to find an English speaker saying something like * “barn + s + yard” in order to pluralize. Instead the whole compound has the plural suffix added to it (“barnyard + s”). Similarly, a prefix is added to the whole compound, not to just the second part; thus we would not hear * “barn + un + yard + like” instead of “un + barnyard + like”. In English there does not seem to be too many restrictions on compounding. Indeed, some commentators refer to this ease of word formation by compounding stems as a very beneficial aspect of English morphology. We may compound noun stems (“fannypack”), noun and verb stems (“strikeforce”), verb stems (“freezedry”), adjective and noun stems (“hardware”), and even a locative and noun (“incrowd”, usually spelled with a hyphen “in-crowd”), and so forth. The reader might try to see how many different kinds of compounds he or she can think of, or even create.

Concord

It should be obvious from the discussion of inflection and derivation that syntax and morphology overlap. In fact it would be quite unproductive to study one without at the same time studying the other. (This is a good opportunity to remind the reader that even though linguistic analysis is presented in this text as though it proceeded step by step from phonetics and phonology to semantics to morphology and finally to syntax, this is definitely not the way linguists work. There may be more emphasis on phonology during the initial stages of analysis, but on the whole all aspects of speech structure are studied simultaneously.) For example, inflectional affixes (morphology) are used to identify syntactic word classes (syntax) while syntactic change is used to identify derivational affixes.

Another area of overlap between morphology and syntax is **concord**. This is the process by which the inflectional affixes of a group of word classes must co-occur, or “be in concord”, with each other. There isn’t much in English to illustrate this, but note that it is similar to the occurrence of the {-Z₃} suffix on a verb in the present tense when the subject is a third person singular (“they/I run” but “she runs”). Spanish, for example, illustrates agreement in many phrases: “this tall boy” /este muchacho alto/ and “this tall girl” /esta muchacha alta/. Here the words in each three word class series (article-noun-adjective) must take certain inflectional suffixes based on the gender suffix of the noun. (These inflectional affixes may be exactly the same, or homophones, as in the Spanish series for “girl”, but this is not always the case, as in the “boy”

series where the inflectional suffix for the article is different (/e/ instead of /o/). Kiswahili is another language which has this kind of process, in fact it has 15 different noun classes, each with its own concord requirements for those adjectives, locatives, articles, and verbal pronominal inflections which are associated with each noun in a sentence (e.g. /kitabu [KI-class noun] kikubwa [adjective stem with ki- concord prefix] kile [demonstrative stem with ki- concord prefix] kiliuzwa [verb stem with ki- concord for proun reference (“it”)]—”Book large that (it) was sold” or, in English, “That large book was sold”). Referring back to the section above on stem and affix identification, it should be obvious that the presence of concord might make affix determination easier since these morphemes tend to stand out due to their repetitive similarities.

Constructions

A word class, as the discussion on agreement should indicate, does not occur in isolation, but rather in more or less close association with other word classes. A group of word classes which has a close association within a phrase or sentence is called a **construction**. The members of a construction are its **constituents**. Identifying constructions is not easy. In some cases the presence of inflectional concord and contiguous placement may assist the linguist. However, the primary method is to elicit informants’ opinions about which word classes can be grouped together as being more closely bonded in contrast to the other word classes in a sentence. A common method is to start with a sentence and ask an informant to divide it into successive halves. For example, serving as my own informant, I would divide “The young boy ran home” into the following two constructions: 1) “The young boy” and 2) “ran home.” Continuing, I would further divide construction 1 into A) “the” and B) “young boy”, and construction 2 into A) “ran” and B) “home;” and finally divide construction 1B into “young” and “boy.” The smallest construction consists of a single word, but larger constructions may consist of many words. If two word classes are constituents in the same construction, then they necessarily have a close syntactic relationship. Thus the adjective word class (in English) has a close syntactic relationship with the noun word class. Constructions whose constituents are an adjective and a noun are very common in English (“young boy,” “large assignment”). But the adjective word class does not have a close syntactic relationship with the verb word class. Try to think of a construction whose constituents are an adjective and a verb.

We saw above that word class membership (which morphemes were members of which word classes) was partly defined by which frames morphemes could occupy, and that another way of defining a sentence is

by a certain sequence of word classes. Using the terms construction and constituent, a “sentence” could be defined as a construction whose constituents were a noun construction and a verb construction (“The boy ran home”), but could not be defined, at least in English, as a construction whose constituents were a locative construction and a verb construction (“*over ran home”). This is the more traditional pedagogical approach to syntax based in semantics which postulates that a grammatical sentence must consist of a sequence of words which “make sense”. Yet it would be interesting to test the proposition that if a speaker said “over ran home” with a /²³¹ ↓ / intonation contour most English hearers would interpret it as an actual “sentence” whose meaning refers to a noun named “over” running home.

Importance of Syntactic Rules in Contemporary Linguistics

While it is quite likely that the sound system of a language is the dominant element of human communication (for example its prominence in early child language acquisition, in morphophonemic variation, or even in word play such as punning and poetry), syntax has come to constitute the major focus of contemporary linguistics. What we have discussed above has taken the orientation of listening carefully to people’s speech and then trying to identify kinds of speech forms and their patterns and relationships (with on-going questioning, of course) and then attempting to work out what mental structures these speakers use to organize their speech. Over the last five to six decades a somewhat different approach has come to the fore, one that attempts to figure out what kinds of cognitive rules must be present in the human mind such that sentences are produced. The goal of linguistics is thus viewed as the delineation of a language’s “grammar” which is the set of the rules necessary to generate all of the well-formed sentences possible in the language, and *only* these sentences (i.e. the set of rules cannot result in an ill-formed sentence. This grammar is usually understood to be vested in human genetic composition (of a part with the evolutionary appearance of *Homo sapiens*) and therefore all humans share the same basic grammar, however different the speech behavior of different human groups might be. Speech is now viewed as a surface condition that is largely secondary to the operations of the deep-seated mental/cognitive grammar. There is a certain appeal to this approach. It tries to gain a much more powerful view of human cognitive processes, in particular in regard to language. It produces very complex, and abstract, generative programs which have a certain similarity to the code used to write computer operating systems. In fact it is closely allied with attempts (still unrealized) to “write” the software for speech production such that an artificial machine or computer (or, often, a robot)

can speak and hold acceptable conversations—the essence of creating artificial intelligence.

It should be no surprise that the descriptions of these grammars result in very complex sets of rules. First, the scope of the grammar has to be extremely large since the number of possible sentences in any language is potentially infinite. Second, speech works with many kinds of conditions and restrictions. This built-in ambiguity and redundancy is an important part of our ability to create complex cultural worlds. For example, a generative grammar for English has to generate “Submarines run quietly” but not “Submarines run home,” while permitting both “Children run quietly” and “Children run home;” transform “Tom reached second base” into “Second base was reached by Tom” but **not** transform the idiomatic “Tom reached for the stars” into “The stars were reached for by Tom;” and allow for the creation of new/unusual, *but acceptable*, sentences such as “John ate the colorless celery” but not create unusual but unacceptable sentences such as **“Colorless celery hibernates in winter.”* These grammars were thus developed to contain many markers and flags in the rules. For example some English nouns are animate and others are not (which is why submarines do not “run home”) and so in the rules nouns had to be marked as being either animate or inanimate. This kind of semantic marking (or flagging) would be extended to all aspects of sentence constituents and the lexicon.

Syntactic organization is perhaps closer to the concerns and interest of speakers (and hearers) than phonological or morphological organization since we typically exchange utterances that are composed of sentences and expressive of whole ideas/commands/desires. Semantic categories therefore serve a much more important role in the construction of generative grammars. The deep mental rules begin with S which represents a thought preparatory to expression in speech (but which has been taken to mean “sentence”). The deep English grammar has the S being replaced by two derivative constructions, a noun phrase (NP) and a verb phrase (VP), the NP is then replaced by (or generated into), in the simplest case, an article and a noun and the VP by a verb and another NP which is in turn replaced by an article and a noun. The grammar could then “look” like this:

S → NP + VP [note that the arrow means “can be replaced by these constituents” and the plus means “and”]

NP → Article + Noun

VP → Verb + NP [A further rule to expand the above NP would not be necessary since these rules are “recursive,” i.e. a speaker would cycle through them over and over again until all constructions are expanded.]

The result would be: Article + Noun + Verb + Article + Noun

This would constitute the basic deep structure. Before becoming speech a lexicon would have to be accessed which would include: Article = “a, the, that....”, Verb = “see, hit....”, and Noun = “boy, girl” And, of course, after items from the lexicon are selected there would have to be a set of phonological rules that would permit the lexical choices to be realized in speech sounds. This grammar would be able to generate such sentences as “The boy saw the girl,” or /²ðə bo¹ sa^w ðə³ gjl¹ ↓ /.

However this set of rules would not be sufficient since, in fact, the verb is a past tense form and the lexicon would supply “see.” Consequently there must be something in the rules that generate the modifications (regular and irregular) in verb form. If S includes the notion of past tense then “saw” would have to be realized, and if the present tense is included in S then “see” would have to be realized. However, if the subject NP is singular there would have to be some rule for adding “s” onto the verb to make it “sees.” Now our set of rules might look like this beginning with the second replacement statements:

NP → Article + Noun + Number

VP → Verb + Tense + NP

and then some rules to take these new elements into account:

Number → singular, plural, and Tense → past, present; with context-specific rules that change the basic structure:

[Article + Noun + singular + Verb + present] →

Article + Noun + Verb + “s” to obtain the verb form “sees.”

(The brackets indicate that the rule only operates on that string of elements—the context by which the rule operates.)

You can begin to appreciate that these generative (or phrase-structure) rules can become quite intricate, even for what would appear to be relatively simple sentences. Furthermore, it is postulated that the active and passive forms of a sentence are based on the same deep structure (which may be true insofar as speakers consider “The boy saw the girl” to

be the same kind of mental entity as “The girl was seen by the boy”). But given this syntactic assumption, the deep rules would have to include some marker as to whether an active or passive sentence would be constructed; if passive then “was” must be inserted (along with the tense) in the VP and “by” must be inserted into the NP, *and* the items must be rearranged (what is termed a transformation rule). Other such generative complexities include the transformations for interrogative constructions also based on the same mental S: “Which boy saw the girl?”, “Where did the boy see the girl?”, and “Which girl did the boy see?”. You may try your hand at constructing a set of generative syntactic rules in exercise 5–7.

That syntactic rules would be very complex is certainly to be expected since language behavior (certainly not only syntax) is the most complex phenomena we know of. To refer back to the example in chapter 1 of two people talking, what we as humans undoubtedly take for granted as a normal part of our interactions, and something that we would not think could be as complex as selecting a college to attend, playing the stock market, or designing an online video game, turns out to be the ultimate challenge for our understanding.

Terms to Know for Chapter Five

morphology	syntax
stem	sentence
affix	word class (“part of speech”)
prefix	— how to define
infix	inflection
suffix	derivation
affix order	root
zero morpheme	compound stems
allomorph	concord
phonologically distributed	construction
allomorphs	constituent
morphologically distributed	generative grammar
allomorphs	transformation rule

APPENDIX A

EXAMPLES OF ARTICULATORY CONFIGURATIONS FROM AUTHOR'S DIALECT

(Western U. S.—Kansas, Oregon, Los Angeles)

Consonants

(All examples given in spelled form.
Dashes indicate sound not in author's dialect.)

<u>articulatory configuration</u>	<u>phonetic graph</u>	<u>initial position</u>	<u>medial position</u>	<u>final position</u>
Simple stops				
bilabial, vl	p	--	spring	lab, gasp
“ , vd	b	big, bring	Hobbit, rabid	lab, scab
alveolar, vl	t	--	British, stop	hat, laughed
“ , vd	d	dark, drink	radish, cardinal	had, poured
palatal, vl	ʃ	--	Ricky, ski	--
“ , vd	ʒ	geese	(Ms) Piggy	--
velar, vl	k	--	Snickers, scum	luck, Hulk
“ , vd	g	gun, grin	sugar	bug, fig
uvular, vl	χ	--		
“ , vd	ʁ	gosh	--	--
Aspirated stops				
bilabial, vl	p ^h	pig, pack	--	--
“ , vd	b ^h	--	--	--
alveolar, vl	t ^h	Tom, time	--	--
“ , vd	d ^h	--	--	--
palatal, vl	ʃ ^h	Kim, quiche	--	--
“ , vd	ʒ ^h	--	--	--
velar, vl	k ^h	come	--	--
“ , vd	g ^h	--	--	--
uvular, vl	χ ^h	cough	--	--

“ , vd	g ^h	--	--	--
Affricated stops				
alveolar, vl	ts	--	Spitzer	watts
“ , vd	d̥z	--	--	(Tiger) Woods
alveopalatal, vl	č	church	pitcher	batch
“ , vd	ǰ	judge	badger, unjust	badge
Fricatives				
labiodental, vl	f	foot, free	puffin	half
“ , vd	v	van	living	live
dental, vl	θ	think, throw	pithy	path
“ , vd	ð	the	feather	lathe
alveolar, vl	s	Sue, swing	assign, east	peace
“ , vd	z	zoo	easy, fizzle	peas
alveopalatal, vl	š	sure, shrink	ocean	mush
“ , vd	ž	Zsa Zsa (Gabor)	azure	rouge
glottal, vl	h	help	--	--
Nasal resonants				
bilabial, vl	m	mark	hammer	ham, I'm
alveolar, vd	n	nice	annual	ban, barn
velar, vd\	ŋ	--	singer	bang
Median resonants				
bilabial, vd	w	wild	--	--
alveolar, vd	r			
alveopalatal, vd	y	yes	--	--
palatal, vd	j	run	torrid	tar
Lateral resonants				
alveolar, vd	l	lead	hilly, early	fill, Earl

Vowels

Simple vowels

Unrounded, nonnasal

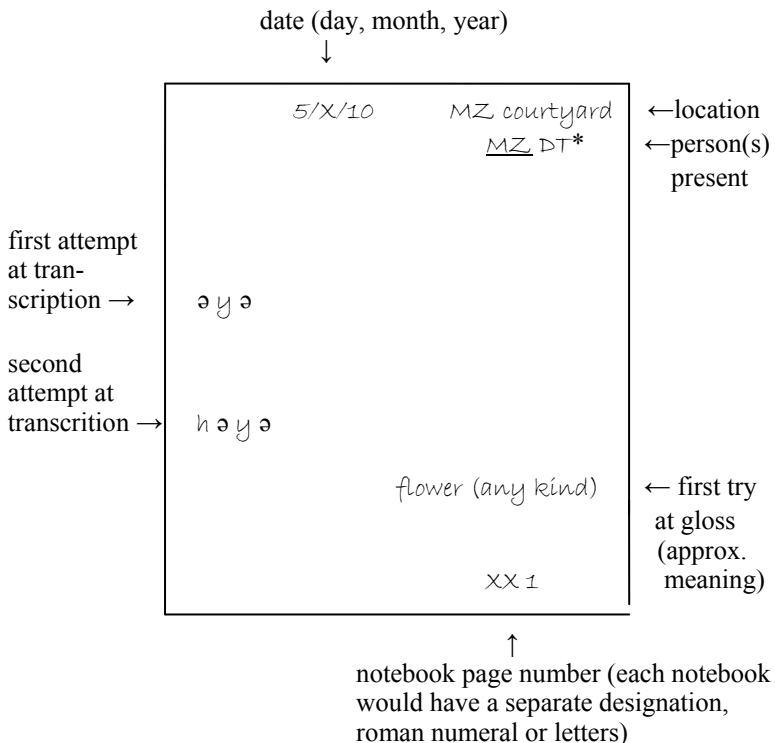
high front	i	it, is	bit, hill	--
mid front	e	egg, x (letter)	bet, gel --	
low front	æ	at, after	bat, shack	--
mid central	ə	ugly, under	but, judge	the
low central	a	all, awful	mall, cause	bah (humbug)
Unrounded, nasal				
high front	i ⁿ	in	sing	--

mid front	ẽ	empty	chemistry	--
(other nasal vowels not shown)				
Rounded, nonnasal				
high back	u	oops	book	--
Glides (complex vowels)				
Forward glides, nonnasal				
from high front unro	i ⁱ	east	sweet	tea
from mid front unro	e ⁱ	ate, aim	bait	bay
from low central unro	a ⁱ	ice	night	lie
from mid back unro	o ⁱ	oyster	void	toy
Back glides, nonnasal				
from low central ro	a ^w	out, ouch	loud	cow
from high back ro	u ^w	ooze	lewd	do, grew
from mid back ro	o ^w	owed, OK	code	tow, beau

APPENDIX B

EXAMPLE OF A PHONETIC CORPUS

I. Example of initial recording (with paper and pen) of speech items (probably words) by asking the names of objects/entities around a household and locality. Each item is recorded on a 3" x 5" spiral-bound notebook page. (Zoque, Mexico, adapted from Gleason, Workbook on Descriptive Linguistics, 1Ed. ©1955 Wadsworth, a part of Cengage Learning, Inc. Reproduced by permission. www.cengage.com/permissions)



6/X/10 MZ field	
	<u>MZ</u>
ko? _ _ _	
ko? _ j i	
ko? n j i	
	bird / fowl chicken turkey
	VI 4

*Only initials or a pseudonym, preferred, is used. The informant is underlined.

II. “Textbook” corpus designed for practice in phonemic analysis. It consists of items taken from notebooks of a field linguist but not organized into any particular classification. For simplicity’s sake no gloss is given, item numbers are for convenience of reference only, and brackets are omitted.

1. həyə	23. nəmgeʔtu	45. tiŋdiŋ
2. haya	24. nəʔs	46. təpkeʔtu
3. kaŋ	25. ɲdʷuxu	47. təʔŋguy
4. kaʔŋji	26. ɲjiŋu	48. tux
5. kənba	27. ɲixpu	49. ʔsəmdʷoʔyu
6. kənduʔyu	28. ŋgama	50. ʔsəmdʷamɲayu
7. keɲaxu	29. ŋgəʔ	51. ʔseht̪su
8. kiʔmdamu	30. ŋgeŋgeʔtu	52. čehčaxu
9. kəʔ	31. ŋgenu	53. čəknaʔču
10. liŋba	32. ŋgyunu	54. tʷətʷəy
11. mbama	33. piŋu	55. winsaʔu
12. mbata	34. petpa	56. wiχtu
13. mbyən	35. poʔk	57. xuʔki
14. minba	36. pyoŋu	58. yaʔsi
15. mindamu	37. pyuŋu	59. ʔaŋjiʔu
16. mingəʔtu	38. pyəŋu	60. ʔanemuč
17. myaŋdamu	39. pəndaʔm	61. ʔaŋdzoŋu
18. ndatah	40. pəŋjəki	62. ʔuy
19. nd̪zima	41. saʔsa	63. ʔyaʔsi
20. nd̪zin	42. suɲi	64. ʔəŋdʷoʔyu
21. ɲjeht̪su	43. šohšaxu	
22. nətʷuxu	44. tatah	

III. Examples of classification of corpus items by selected phonetic criteria. Each item is listed along with approximate gloss. Item number if notebook number and page number.

bilabial initials—resonants

I 16 [mbama]	“my clothing”
II 18 [mbata]	“my mat”
III 1 [mbyən]	“you’re a man”
III 5 [minba]	“he comes”
III 20 [mindama]	“you came”
III 21 [mingɛʔtu]	“he also came”
III 22 [myaŋdamu]	“you went”

bilabial initials—stops

I 40 [piŋu]	“he picked it up”
I 46 [petpa]	“he sweeps”
I 52 [poʔk]	“knot”
II 3 [pyoŋu]	“he burned it”
II 15 [pyuŋu]	“he scattered it”
II 17 [pyəŋu]	“he broke it”
IV 5 [pəndaʔm]	“men”
IV 6 [pəŋʂəki]	“image of a man”

APPENDIX C

THE INTERNATIONAL PHONETIC ALPHABET (2005)

CONSONANTS (PULMONIC)

	Bilabial	Labio-dental	Dental	Alveolar	Post-alveolar	Retroflex	Palatal	Velar	Uvular	Pharyngeal	Epi-glottal	Glottal
Nasal	m	ɱ		n		ɳ	ɲ	ŋ	ɴ		ʔ	ʕ
Plosive	p b	ɸ β	t d			ʈ ɖ	c ɟ	k ɡ	q ɢ		ʔ	ʕ
Fricative	ɸ β	f v	θ ð	s z	ʃ ʒ	ʂ ʐ	ç ʝ	x ɣ	χ ʁ	ħ ʕ	ʜ	ɦ
Approximant		ʋ		ɹ		ɻ	j	ɰ			ʕ	ɦ
Trill	ʙ		ʀ						ʀ		ʕ	
Tap, Flap		ⱱ	ɾ	ɽ								
Lateral fricative			ɬ ɮ	ɮ								
Lateral approximant			ɭ	ɭ			ʎ	ʎ				
Lateral flap			ɺ	ɺ								

Where symbols appear in pairs, the one to the right represents a modally voiced consonant, except for murmured *f*.
Shaded areas denote articulations judged to be impossible. Light grey letters are unofficial extensions of the IPA.

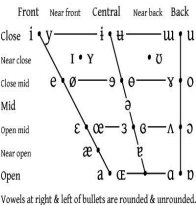
CONSONANTS (NON-PULMONIC)

Anterior click releases (require posterior stops)	Voiced implosives	Ejectives
ɔ Bilabial fricated	ɓ Bilabial	ʼ Examples:
ɬ Laminar alveolar fricated ("dental")	ɗ Dental or alveolar	ɓ Bilabial
ɮ Apical (postalveolar) abrupt ("retroflex")	ɗ Palatal	ɗ Dental or alveolar
ɮ Laminar postalveolar abrupt ("palatal")	ɗ Velar	ɗ Velar
ɮ Lateral alveolar fricated ("lateral")	ɗ Uvular	ɗ Alveolar fricative

CONSONANTS (CO-ARTICULATED)

- M Voiceless labialized velar approximant
- W Voiced labialized velar approximant
- ɥ Voiced labialized palatal approximant
- ɕ Voiceless palatalized postalveolar (alveolo-palatal) fricative
- ʑ Voiced palatalized postalveolar (alveolo-palatal) fricative
- ɥ Simultaneous x and j (disputed)
- kp ts Affricates and double articulations may be joined by a tie bar

VOWELS



SUPRASEGMENTALS

- ˈ Primary stress
- ˌ Secondary stress
- ː Long
- ˑ Short
- ˑ Extra-short
- ˑ Syllable break
- ˑ Linking (no break)
- ˑ Minor (foot) break
- ˑ Major (intonation) break
- ˑ Global rise
- ˑ Global fall
- ˑ Level tones
- ˑ Contour-tone examples
- ˑ Top
- ˑ High
- ˑ Mid
- ˑ Low
- ˑ Bottom
- ˑ Tone terracing
- ˑ Upstep
- ˑ Downstep
- ˑ Rising
- ˑ Falling
- ˑ High rising
- ˑ Low rising
- ˑ High falling
- ˑ Low falling
- ˑ Peaking
- ˑ Dipping

DIACRITICS

Diacritics may be placed above a symbol with a descender, as *ɰ*. Other IPA symbols may appear as diacritics to represent phonetic detail: *ɾ* (fricative release), *ɸ* (breathy voice), *ʔ* (glottal onset), *ʕ* (epenthetic schwa), *ʔ* (diphthongization).

SYLLABILITY & RELEASES		PRONATION		PRIMARY ARTICULATION		SECONDARY ARTICULATION			
ɲ ɳ	Syllabic	ɲ ɳ	Voiceless or Slack voice	t̪ d̪	Dental	tʷ dʷ	Labialized	ɰ ɰ	More rounded
ɸ ɹ	Non-syllabic	ɸ ɹ	Marginal voice or Shift voice	t̪ʰ d̪ʰ	Apical	t̪ʰ d̪ʰ	Palatalized	ɰʷ ɰʷ	Less rounded
t̪ʰ t̪ʰ	(Pre)aspirated	ɲ̤ ɲ̤	Breathy voice	t̪ d̪	Laminal	t̪ʰ d̪ʰ	Velarized	ẽ ẽ	Nasalized
d̪ⁿ	Nasal release	ɲ̤ ɲ̤	Creaky voice	ɰ t̪	Advanced	t̪ʰ d̪ʰ	Pharyngealized	ɰ ɰ	Rhoticity
ɰⁿ	Lateral release	ɲ̤ ɲ̤	Strident	t̪ʰ d̪ʰ	Retracted	t̪ʰ d̪ʰ	Velarized or pharyngealized	ɰ ɰ	Advanced tongue root
t̪	No audible release	ɲ̤ ɲ̤	Lingual	ɰ	Centralized	ɰ	Mid-centralized	ɰ ɰ	Retracted tongue root
ɸ ɹ	Lowered (ɸ is a bilabial approximant)	ɲ̤ ɲ̤		ɰ	Raised (ɰ is a voiced alveolar non-sibilant fricative)				

APPENDIX D

ENGLISH INFLECTIONAL AND DERIVATIONAL AFFIXES

Stem Word Class	Inflectional Affixes	Derivational Affixes	Changes Stem Into
noun	{-Z ₂ } → /-s/ “book <u>s</u> ” /-z/ “ball <u>s</u> ” /- z/ “rose <u>s</u> ”	{-Z} “-’s (possessive)” “the book’ <u>s</u> content”	adjective
	/səb-/ “sub-”, “ <u>sub</u> class”	{-IN} “an/-ian” “American”	noun adjective
	/nan-/ “non-”, “ <u>non</u> entity”		
adjective	/æntə ⁱ -/ “anti-”, “ <u>anti</u> hero”	/- əl/ “-al”, “national <u>l</u> ” (+V change in stem)	
	{VM-} →/ n-/ “ <u>un</u> happiness” /im-/ “ <u>im</u> propriety” /in-/ “ <u>in</u> hospitality”	/-li ⁱ /, /-i ⁱ / “-ly”, “-y” “man <u>ly</u> ”, “ros <u>y</u> ”	
adjective		/-ful/ “-ful”, “thought <u>ful</u> ”	
adjective		/-iʃ/ “ish”, “boy <u>ish</u> ”	
		/-ik/ “-ic” “poet <u>ic</u> ” + stress change in stem)	adjective
		/-les/ “-less”, “home <u>less</u> ”	adjective
		/-izm/ “-ism”, “sex <u>ism</u> ”	noun

		/-æj ⁱ / “-ary”, “caution <u>ary</u> ” adjective
		/en-/ “ <u>enslave</u> ” verb
		/ænta ⁱ -/ “anti-”, “ <u>antitank</u> ” adjective
verb	{-Z ₃ } “-s”, “he looks <u>s</u> ”	/-ej/ “-er”, “runner <u>s</u> ” noun
	{-D ₁ } →/-ed/ “planted”	{-D ₂ } “a pleased <u>s</u> customer” adjective
	/-d/ “ <u>rigged</u> ”	/-iŋ/ “jogging is a hobby”
noun	/-t/ “danced <u>s</u> ”	“ “jogging shoes” adjective
	/mis-/ “mis-”, “ <u>mis</u> understand”	/-əbəl/ “this is a doable --” adjective
	/ji ⁱ -/ “re-”, “ <u>re</u> think it”	
	/ən-/ “un-”, “ <u>un</u> do the knot”	/-ins/ “-ence”, “emerg <u>ence</u> ” noun
		/-ment/ “-ment”, “enjoy <u>ment</u> ” noun
adjective		/-fəl/ “-full”, “wakeful”
		/-ʃən/ “-tion”, “ <u>action</u> ” noun
adjective		/-nes/ “wakeful <u>ness</u> ” noun
	/-ej/ “-er”, “whit <u>er</u> ”	/-li ⁱ / “-ly”, “slow <u>ly</u> ” adverb
	/-ist/ “-est”, “whit <u>est</u> ”	/-izm/ “-ism”, “national <u>ism</u> ” (+V change in stem) noun
	/mis-/ “mis-”, “ <u>mis</u> taken” (with stem = verb + {-D ₂ })	/-a ⁱ z/ “-ize”, “nationalize” verb
verb	/nan-/ “non-”, “ <u>non</u> allergic”	/en-/ “en-”, “ <u>en</u> dear”
verb	{VM-} →/im-/ “ <u>im</u> proper”	/-en/ “-en”, “shar <u>pen</u> ”
	/in-/ “ <u>in</u> hospitable”	
	/ən-/ “ <u>un</u> happy”	

/əltʃə-/ “ultra”, “ <u>ultra</u> smooth”	
/jiː-/ “re-”, “ <u>re</u> worked”	
(with stem = verb + D ₂)	

BIBLIOGRAPHY

- Abercrombie, David. 1967. *Elements of General Phonetics*. Chicago: Aldine.
- Ashby, Patricia. 1995. *Speech Sounds*. London: Routledge.
- Atkinson, Martin, D. Kilby, and I. Roca. 1982. *Foundations of General Linguistics*. London: Allen & Unwin.
- Carroll, Lewis. 1960 [1871]. *In The Annotated Alice. Alice's Adventures in Wonderland & Through the Looking Glass*. Martin Gardner, ed. New York: Bramhall House.
- Cherry, Colin. 1961. *On Human Communication: A Review, a Survey, and a Criticism*. New York: Science Editions, Inc.
- Crowley, Terry. 2007. *Field Linguistics: A Beginner's Guide*. New York: Oxford University Press.
- Crystal, David. 1971. *Linguistics*. Hammondsworth, UK: Penguin.
- . 1997. *The Cambridge Encyclopedia of Language*, 2nd Edition. Cambridge: Cambridge University Press.
- . 2009. *Just a Phrase I'm Going Through: My Life in Language*. New York: Routledge.
- . 2010. *A Little Book of Language*. New Haven: Yale University Press.
- Denes, Peter B., and Elliot N. Pinson. 1963. *The Speech Chain: The Physics and Biology of Spoken Language*. New York: Bell telephone Laboratories.
- Everett, Daniel L. 2012. *Language: The Cultural Tool*. New York: Pantheon Books.
- Evans-Pritchard, E. E. 1940. *The Nuer*. London: Oxford University Press.
- Finnegan, Ruth. 2002. *Communicating: Multiple Modes of Human Interconnection*. New York: Routledge.
- Flexner, Stuart Berg. 1976. *I Hear America Talking: An Illustrated History of American Words and Phrases*. New York: Touchstone.
- Gaeng, Paul A. 1971. *Introduction to the Principles of Language*. Lanham, MA: University Press of America.
- Gleason, Henry A. 1955. *Workbook in Descriptive Linguistics*. New York: Holt, Rinehart, and Winston.
- . 1961. *An Introduction to Descriptive Linguistics*, Rev. Ed. New York: Holt, Rinehart, and Winston.
- Greenberg, Joseph. 1977. *A New Invitation to Linguistics*. Garden City: Anchor Press.

- Groves, Colin. 2006. Primate Taxonomy. *In* Encyclopedia of Anthropology. H. James Bix, ed. Thousand Oaks, California: Sage Publications.
- Hockett, Charles. 1958. A Course in Modern Linguistics. New York: Macmillan.
- Howell, P. P. 1954. A Manual of Nuer Law. London: Oxford University Press.
- Hymes, Dell (Ed.). 1964. Language in Culture and Society: A Reader in Linguistics and Anthropology. New York: Harper & Row.
- Lehmann, Wilfred. 1962. Historical Linguistics: An Introduction. New York: Holt, Rinehart, and Winston.
- Maddieson, Ian. 1984. Patterns of Sounds. Cambridge, UK: Cambridge University Press.
- O'Grady, William, M. Dobrovolsky, and M. Aronoff. 1989. Contemporary Linguistics. New York: St. Martin's Press.
- Okrand, Marc. 1985. The Klingon dictionary: English/Klingon, Klingon/English. New York: Pocket Books.
- Poirier, Frank. 1987. Understanding Human Evolution. Englewood Cliffs, NJ: Prentice-Hall.
- Pullum, G. K. and W. A. Ladusaw. 1986. Phonetic Symbol Guide. Chicago: University of Chicago Press.
- Samarin, William J. 1967. Field Linguistics. New York: Holt, Rinehart, and Winston.
- Samson, Geoffrey. 1980. Schools of Linguistics. Stanford, CA: Stanford University Press.
- Trudgill, Peter. 1983. Sociolinguistics, 2nd Ed. Hammondsorth, UK: Penguin.
- Witucki, Jeannette. 1984a. Introducing Linguistic Analysis: Phonemics. South Pasadena, CA: Condor Book Company.
- . 1984b. Introducing Linguistic Analysis: Morphology and Syntax. South Pasadena, CA: Condor Book Company.
- Whorf, Benjamin L. (1940). Linguistics as an Exact Science. IN John Carroll (Ed.), Selected Writings of Benjamin Lee Whorf. Cambridge, MA: M.I.T. Press 1956.

EXERCISES FOR CHAPTERS TWO–FIVE

(suggested solutions in following section)

Exercises For Chapter Two: Phonetics

Guidelines

Each of the following exercises will ask you (by yourself or with others) to say words or phrases and describe the sounds which make up these samples of speech. However there are several important guidelines which you should follow.

1) Say the words or phrases in a normal manner. Don't speak so slowly that you over-emphasize each syllable, nor so rapidly that syllables or even words in a phrase are run together or distorted by being shortened. For example, exercise number 1 asks you to say your name. You should say it just as you would if you were saying it in a phrase such as, "Hello, my name is _____." This whole phrase should last a little under two seconds if you just use your first name, one second is too fast and four or more seconds is too slow. This will give you a guide to how fast you should say just the name. In fact a good strategy is to say the word or phrase as part of some larger normal speech context, just to establish a normal pronunciation, and then just say the individual word or phrase at that speed as many times as you need to. It will likely be necessary for you to hold your articulatory motions at certain points in order to better sense the nature of the articulatory configuration at these points. But always return to a normal speed when you check your description of the sounds.

2) You will also have to write down these words and phrases using their spelled forms before you start to describe them in articulatory terms. When you are saying the words while trying to figure out how each sound is articulated, **DO NOT LOOK AT THE SPELLED VERSION**. The spelling will only confuse you since there will rarely be a one-to-one correspondence between letters and sounds. (For example, "gate" has three sounds, not four.) Further, the letters themselves usually do not have the same articulatory significance as we are using here.

3) Finally, do not feel that you must get the exact articulatory features of a sound every time on the first attempt. If you are unsure about a sound then identify which alternatives you think it may be, for example you may think a consonantal sound is either palatal or velar, but clearly not alveolar or uvular. While you may not have precisely identified it you have

achieved a narrowing of the possibilities. In fact, a system of successive approximation is what linguists themselves do. As you gain experience you will find that you will be able to make closer and closer estimates of difficult or unfamiliar sounds.

The examples given for the exercises are based upon the author's speech and are consequently meant as illustrations, not as approved or correct renditions. Your speech forms may vary from these illustrations, consequently your task is to describe the speech you are observing (yours or that of a friend), not to copy the author's.

Example for Exercises 2–1 to 2–3

<u>Place of Articulation</u>			
<u>“Mariana”</u>	<u>Manner</u>	<u>Consonants</u>	<u>Vowels</u>
M	resonant	bilabial	
a	vowel		mid front, or low front
r	resonant	alveo-palatal or palatal	
i	vowel		high front
a	vowel		low central, or low front
n	resonant	alveolar	
a	vowel		mid central

<u>Place of Articulation</u>			
<u>“John”</u>	<u>Manner</u>	<u>Consonants</u>	<u>Vowels</u>
J	stop and then fricative	alveolar, or alveo- palatal	
o \ (one h / sound)	vowel		low central, or low back
n	resonant	alveolar	

Note: The letters of “Mariana” and “John” are listed alongside the phonetic descriptions of the actual speech sounds in both. This is just an initial aid, but remember you **MUST LISTEN to the sounds and not look at the letters**, so even if you write down a word or are given a word to describe, don't look at the letters. Note that while for “Mariana” each letter corresponds to an articulatory event, in “John” the two middle letters actually correspond to one.

2–1 Say your name (or your nickname). The exercise is for you to list, in sequence, the **manners** of articulation that you used for the sounds that make up your name. (See the example above in regard to manner.)

(YOUR NAME—spell on line:)

Sequence of
Manners of articulation

(spell name downward)↓

2–2 Say your name, the same one you used for Exercise 1. Now list, in sequence, the **places** of articulation that you used for these sounds. (See the example above in regard to place.)

(YOUR NAME—spell on line:)

Sequence of
Places of Articulation
(just for consonants)

(spell name downward)↓

2–3 Again, say your name (as for exercises 2–1 and 2). Now list, in sequence, the kinds of vowels that you used (omit the consonants)

(YOUR NAME)

Sequence of Kinds of Vowels

(spell name downwards)↓

2–4 Write down the names of four objects or places in your room or home (spell them as you would normally do). Format the names as done in the examples above, leaving space to write in the manner and place of articulation of each sound (not necessarily each letter) with separate columns for consonant and vowel

(Here is an example)

“bike”

lettersmannerplace for
consonantsplace for
vowels

b

stop

bilabial

i

vowel

low central (?and
then high front)

k\ [only stop

velar

e/ one sound]

2–5 Using any **five** from one of the sets of ten words given below, list the complete manner and place of articulation for each consonant sound. (**Note that the format for your answer is changed** from the format for Exercises 1–4. Look at the example.)

Example:

“desk”

“d

e

s

k”

manner: stop

V

fricative

stop

voicing: vd

V

vl

vl

aspiration: unasp

V

unasp

unasp

place: alveolar

V

alveolar

velar

[Note: vd = voiced, vl = voiceless, asp = aspirated, unasp = unaspirated]

List A: “watch, lamp, glove, spoon, shelf, close, thumb, shirt, slip, scrub”

List B: “sink, fruit, bunch, vase, bat, mail, hot, judge, sell, mouse”

2–6 This is the same kind of exercise as 2–5, but here you are to list a complete description of the vowels. You should use the same words you used for Exercise 5.

Example “rayon”

	“r	a	y	o	n”
complex/simple:	C	forward glide	C	simple	C
height:	C	mid —>high	C	low	C
position:	C	front	C	back	C
rounding:	C	unr	C	unr	C
nasalization:	C	nonn	C	nonn	C
length:	C	---	C	---	C

[Note: unr = unrounded, nonn = nonnasal]

2–7 This exercise will sharpen your definitional skills on the articulatory terms that have been presented for consonant and vowel description. Briefly, but as accurately as you can, explain the significant articulatory differences between each pair of terms:

stop—fricative	velar stop—glottal stop
alveolar—alveo-palatal	bilabial stop—bilabial nasal resonant
median resonant—lateral resonant	voiced C—voiceless C
aspirated stop—affricate	high V—low V
front V—back V	complex V (glide)—simple V
forward glide—back glide	rounded V—unrounded V
resonant—vowel	nasalized V—nonnasalized V
bilabial—labiodental	

2–8A This exercise gives you practice in using the system of phonetic notation presented in the chapter. Using the system of phonetic notation presented in Figure 2.8A, represent the consonants in the words of one of the columns below. Just write V where there is a vowel.

(Example “Anthro” would be [V n θ j V])

<u>A</u>	<u>B</u>	<u>C</u>
bank	beauty	pure
city	ghost	passion
fine	cheer	schedule
time	cork	courses
Dave	once	knot
three	sharp	enough

2–8B Now, using the same words as for 8A, represent each of the vowels using the notation shown in Figure 2.8B. Just write C where there is a consonant. (Example “Anthro” would be [æⁿ C C C o^w])

2–9 For any two of the polysyllabic words listed below, identify which syllable has the primary stress, and, if there are three or more syllables, which syllable has the secondary stress. The words are spelled so you should transcribe them into phonetic notation first. Each vowel marks a syllable; place your stress graphs over the appropriate vowels.

(Example: “double” [d ə b ə l], “anthropology” [ʌⁿ n θ j o^w p á l o^w ʃ iⁱ]
language primary dental linguistics articulation phonetics)

2–10 This last exercise is to be done with another person studying linguistics. Select some short segment of speech such as two lines from a movie, a slogan, a poem, or even the first two line of your club’s mission statement. Do not tell the other person what you have selected. Represent it in phonetic notation and then give it to your partner, collecting his or hers. Can each of you say what the other has written? If not, what were their shortcomings?

Exercises for Chapter Three

3–1 For this exercise you are to use your own speech as the basis for finding minimal pairs of selected English speech sounds. You should try to find minimal pairs in initial, medial, and final positions for each pair of sounds.

If you do not have a minimal pair in one of the positions, then write “none” in the space. List ONE pair in each space, represent the word in English spelling AND by phonetic notation (using phonetic graphs). One example is given:

<u>phoneme pair</u>	<u>initial</u>	<u>medial</u>	<u>final</u>
[p] [b]	none	“rapid” [jæpid]	“lip” [lip]
		“rabid” [jæbid]	“lib” [lib]

[f] [v]

[f] [θ]

[s] [ʃ]

[s] [z]

[k] [g]

[n] [ŋ]

[i] [æ]

[e] [a]

3–2 This exercise is to gain practice in constructing a phonetic chart of the sounds in a corpus. This is the first step in doing a phonemic analysis of the corpus, especially for a distributional analysis. Using the Zoque (Mexico) corpus presented in Appendix B, page 129, fill in the following vowel charts (consonants will not be done for this exercise): (Adapted from Gleason, *Workbook in Descriptive Linguistics*, 1Ed. ©1955:65 Wadsworth, a part of Cengage Learning, Inc. Reproduced by permission. www.cengage.com/permissions)

Zoque Vowels

		UNROUNDED Location of Tongue Height						
		Front		Central		Back		
		nasal	non-nasal	nasal	non-nasal	nasal	non-nasal	
Tongue Height	High							
	Lower-high							
	Mid							
	Lower-mid							
	Low							

		ROUNDED Location of Tongue Height						
		Front		Central		Back		
		nasal	non-nasal	nasal	non-nasal	nasal	non-nasal	
Tongue Height	High							
	Lower-high							
	Mid							
	Lower-mid							
	Low							

3–3 This exercise is to gain practice in constructing a phonetic chart of the sounds in a corpus. This is the first step in doing a phonemic analysis of

the corpus, especially for a distributional analysis. Using the following corpus of Chamorro utterances (Guam, after Witucki 1984a:137), construct a phonetic chart of the consonants present in the corpus.

- | | | | |
|--------------|---------------|----------------|---------------|
| 1. [məsgoʔ] | “juicy” | 14. [bɛʔdi] | “green” |
| 2. [pæʔgon] | “child” | 15. [si:ha] | “they” |
| 3. [pokpok] | “swelling” | 16. [sɒmnæk] | “sun” |
| 4. [mæmpɒs] | “very” | 17. [tæ: noʔ] | “earth” |
| 5. [ʔa:ga] | “ripe banana” | 18. [tæ:si] | “sea” |
| 6. [hi:ta] | “we” | 19. [foʔgon] | “wet” |
| 7. [ʔu:loʔ] | “worm” | 20. [læ: ŋi t] | “sky” |
| 8. [ʔæ:su] | “smoke” | 21. [fugoʔ] | “to squeeze” |
| 9. [hu:læʔ] | “tongue” | 22. [ni:fiɪm] | “tooth” |
| 10. [gu:mæʔ] | “house” | 23. [di:doʔ] | “sharp” |
| 11. [ʔæŋgm] | “if” | 24. [bu:la] | “full derous” |
| 12. [peknoʔ] | “murderous” | 25. [bokboʔ] | “to pull out” |
| 13. [hæ:lɒm] | “inside” | | |

[V:] indicates a long vowel, [ʊ] is a lower-high back rounded, non-nas. vowel

Chamorro Consonants

	bi-labial		labio-dental		dental		alveolar		alveo-palatal		palatal		velar		uvular		glottal
	vl	vd	vl	vd	vl	vd	vl	vd	vl	vd	vl	vd	vl	vd	vl	vd	vl
simple stops																	
aspirated stops																	
affricated stops																	
fricatives																	
nasal resonants																	
lateral resonants																	
median resonants																	

3–4 Use the method of distributional analysis to determine the phonemic status of the consonants [β] and [b] in the following corpus. Spanish (Castillian, after Witucki 1984a:64)

- | | | | |
|----------------|-------------|--------------------|------------|
| 1. [laβamanos] | “washstand” | 6. [leβaðo] | “carried” |
| 2. [bonita] | “pretty” | 7. [bai iiiiɪ laʃ] | “to dance” |
| 3. [loβo] | “wolf” | 8. [boka] | “mouth” |
| 4. [liβɛʔtað] | “liberty” | 9. [biβo] | “alive” |

5. [uniβeʁʂiðað “university” 10. [bəʁβo] “verb”

Note: [β] and [ð] are voiced fricatives (bilabial and dental respectively) and ʁ is an alveolar flap (tongue tip quickly strikes the alveolar ridge and pulls away).

Are these two sounds in complementary distribution?

Describe the distribution:

What is their phonemic status?

3–5 Use the method of distributional analysis to determine the phonemic status of the consonants [ɣ] and [x] in the following corpus. Chiricahua Apache (southwestern U. S.; after Witucki 1964a:63)

- | | | | |
|---------------|------------------|-----------------|---------------------|
| 1. [ɣe:ɫ] | “pack” | 5. [ʂisɣeh] | “I kill him” |
| 2. [xoʂ] | “cactus” | 6. [deɫɣiz] | “I have twirled it” |
| 3. [xaʔ] | “winter” | 7. [xaʔya] | “where” |
| 4. [ʔiɫxoʂ] | “he is sleeping” | 8. [ha:stɣi:] | “old man” |

Note: V: is a long vowel, [ɣ] is a voiceless palatal fricative, [x] is a voiceless velar fricative, and [ɫ] is a voiceless alveolar lateral resonant

Are these two sounds in complementary distribution?

Describe the distribution:

What is their phonemic status?

3–6 Use the method of distributional analysis to determine the phonemic status of the Kiswahili vowels [o] and [ɔ] in the following corpus Kiswahili (East Africa) (Adapted from Gleason, *Workbook in Descriptive Linguistics*, 1Ed. ©1955:58 Wadsworth, a part of Cengage Learning, Inc. Reproduced by permission. www.cengage.com/permissions)

- | | | | |
|---------------|------------|----------------|--------------|
| 1. [ŋɔma] | “drum” | 6. [ʝogo] | “rooster” |
| 2. [boma] | “fort” | 7. [ñɔña] | “nurse” |
| 3. [kuona] | “to see” | 8. [kuokota] | “to pick up” |
| 4. [ndoto] | “a dream” | 9. [kʰɔndo] | “sheep” |
| 5. [watoto] | “children” | 10. [moja] | “one” |

Note: [ɔ] is a rounded low back non-nasal vowel; [ŋ] is an alveo-palatal nasal resonant

Are these two sounds in complementary distribution?

Describe the distribution:

What is their phonemic status?

3–7 Use the method of distributional analysis to determine the phonemic status of the consonants [n] and [m] in the following corpus Chamorro (Guam, Mariana Islands—Western Pacific; after Witucki 1984a:46)

- | | |
|----------------------------------|------------------------------|
| 1. [nunu] “banyan tree” | 6. [mamis] “sweet” |
| 2. [nidzok] “mature coconut” | 7. [hanom] “water” |
| 3. [nikə] “wild yam” | 8. [amot] “medicine” |
| 4. [hanon] “to burn something” | 9. [mumu] “to fight” |
| 5. [nana] “mother” | 10. [midzʊ] “you (plural)” |

Note: [ʊ] is a lower-high back rounded vowel, [ɪ] is a lower-high, front unrounded vowel

Are these two sounds in complementary distribution?

Describe the distribution:

What is their phonemic status?

3–8 Use the method of distributional analysis to determine the phonemic status of the consonants [k] and [g] in the following corpus Manx (Isle of Man—United Kingdom; after Witucki 1984a:47)

- | | |
|-------------------------------|------------------------------|
| 1. [mak] “son” | 8. [ɡetlax] “flying” |
| 2. [kənax] “buying” | 9. [bɛɡ] “beside” |
| 3. [ket ^h] “cat” | 10. [ɡærəs] “being hungry” |
| 4. [krɛk] “selling” | 11. [ɡɛnəl] “merry” |
| 5. [klag] “clock” | 12. [grɛ] “saying” |
| 6. [kælax] “cock” | 13. [ɡol] “going” |
| 7. [mənɪk] “frequent” | 14. [momɪɡ] “mother” |

NOTE: [ɛ] is a lower-mid front unrounded vowel, [ɪ] is a lower-high front unrounded vowel

Are these two sounds in complementary distribution?

Describe the distribution:

What is their phonemic status?

3–9 Refer to the Zoque (Mexico) corpus in Appendix B (page 129) and exercise 3–2 above. Using the method of distributional analysis determine the phonemic status of the stop consonants [p] and [b]. Referring to the principle of allophonic consistency, predict the phonemic status of the stops [t] and [d], [tʲ] and [dʲ] (these are alveo-palatal stops), and [k] and [g]. Test your prediction and write your conclusions at the bottom. Do a distributional analysis for these two phones: [p] and [b] make a prediction about the phonemic relationship of Zoque voiced and voiceless stops based upon your analysis.

Now test your prediction by doing a distributional analysis of one of the pairs below:

[t] and [d]

[tʸ] and [dʸ]

[k] and [g]

What is your prediction concerning the general principle of the phonemic status of voicing in Zoque?

3–10 Use the Zoque (Mexico) corpus in Appendix B (page 129). Assume that a complete phonemic analysis has determined that Zoque stop phonemes are as follows: Each pair of phones represents one phoneme—i.e. there are 6 phonemes, each with two allophones—[p, b], [t, d], [tʸ, dʸ], [t̥, d̥], [č, ǰ], and [k, g].

Decide upon a system of phonemic graphs for each of these phonemes. What is the basis for your selection?

Fill in the phoneme graphs in the slant lines for each group of allophones:
 / / [p, b]; / / [t, d]; / / [tʸ, dʸ]; / / [t̥, d̥]; / / [č, ǰ];
 / / [k, g]

Now rewrite phonemically each of the following Zoque items (by their item number). You can assume that the other consonant and vowel phoneme graphs are the same as their phonetic graphs.

Rewrite item numbers: 3; 4; 12; 19; 21; 29; 34; 45; 49; and 52.

Exercises for Chapter Four

Example of morpheme identification exercises:

Here is a small corpus from the Mende language (West Africa, Sierra Leone) (Adapted from Gleason, *Workbook in Descriptive Linguistics*, 1Ed. ©1955:23 Wadsworth, a part of Cengage Learning, Inc. Reproduced by permission. www.cengage.com/permissions)

For convenience sake, it is organized into two columns and a comparison of these two columns will demonstrate the presence of a fifth morpheme.

- | | |
|--------------------|------------------------|
| 1. /pélé/ “house” | 5. /pélé/ “the house” |
| 2. /káli/ “hoe” | 6. /káli/ “the hoe” |
| 3. /návo/ “boy” | 7. /návo/ “the boy” |
| 4. /númú/ “person” | 8. /númú/ “the person” |

List all the morphemes you can identify and give their approximate meanings.

The task is to identify as many of the morphemes (those sequences of phonemes which co-occur with a gloss element/meaning) as you can from the corpus. You should see that there is evidence for five morphemes. Four of these stand alone in the left-hand column (do not think that this

will always be the case, or that any of the morphemes in a corpus will stand alone—be free morphemes). You need to check that the phoneme strings/sequences /pélé/, /káli/, /návó/, and /númú/ co-occur with (respectively) “house”, “hoe”, “boy”, and “person” in both columns. Which is the fifth morpheme? You should see that there is one other phoneme string not yet identified and that it co-occurs with the gloss “the”: /i/. Do not be put off by the fact that this is only one phoneme, many morphemes are only one phoneme. This is the fifth morpheme, and then this is how your answer would look: /pélé/ “house” /káli/ “hoe” /návó/ “boy” /númú/ “person” /i/ “the”.

You are asked for an approximate meaning even though providing a suitable gloss in this exercise does not seem too difficult. But this will often be a real challenge. Try to use the actual gloss in the corpus as much as you can. This may not be the most accurate, but it is certainly easier than figuring out an embedded meaning—although that is exactly what you will have to do in most cases. Simply write the gloss in quotation marks, as was done here. However, if you must use an embedded meaning or one you deduce from the glosses, these must be placed in square brackets inside the quotation marks. For example, if you believe that the initial /n/ in items 3 and 4 might be a morpheme indicating an animate entity (“boy” and “person” as opposed to “house” and “hoe”) then you would have to include in your answer: /n/ “[animate entity or living thing]”. Keep in mind that if you identify /n/ as a morpheme **then** you will have to describe the remaining phoneme strings (/ávó/ and /úmú/) as morphemes too and come up with glosses for them as well. ALL of the phonemes must be accounted for as part of a morpheme—there are no “free” or unrelated phonemes.

4-1 Classical Nahuatl (Central Mexico, Aztec, after Witucki 1984b:9).

List all the morphemes and give their approximate meanings

- | | | | |
|---------------|-------------|-----------------|--------------|
| 1. / notew / | “my stone” | 7. / iitew / | “his stone” |
| 2. / nopetl / | “my mat” | 8. / iipetl / | “his mat” |
| 3. / nota? / | “my father” | 9. / iita? / | “his father” |
| 4. / nokwik / | “my song” | 10. / iikwik / | “his song” |
| 5. / notlew / | “my fire” | 11. / iitlew / | “his fire” |
| 6. / noamaw / | “my book” | 12. / iiaamaw / | “his book” |

NOTE: double vowel graphs indicate a long vowel (e.g. /ii/ = a long /i/), /l/ = a voiceless alveolar lateral resonant

List all morphemes and give their approximate meanings:

4–2 Delaware (Amerind, eastern woodlands, after Witucki 1984b:7)

- | | | | |
|-----------------|----------------------|---------------|-----------------|
| 1. / nuuxtət / | “my little father” | 7. / ktaan / | “your daughter” |
| 2. / kuuxtət / | “your little father” | 8. / ntaan / | “my daughter” |
| 3. / wuuxtət / | “his little father” | 9. / wtaan / | “his daughter” |
| 4. / wsiittət / | “his little foot” | 10. / nuux / | “my father” |
| 5. / ntuuntət / | “my little mouth” | 11. / wsiit / | “his foot” |
| 6. / ktuuntət / | “your little mouth” | | |

Note: a double vowel graph indicates a long vowel (e.g. /uu/ = long /u/), however a double consonant graph actually indicates two successive consonants.

List all the morphemes and give their approximate meanings:

4–3 Kiswahili (East Africa) (Adapted from Gleason, *Workbook in Descriptive Linguistics*, 1Ed. ©1955:26 Wadsworth, a part of Cengage Learning, Inc. Reproduced by permission. www.cengage.com/permissions)

- | | | |
|-----|-----------------|------------------------|
| 1. | / atanipenda / | “he will like me” |
| 2. | / atakupenda / | “he will like you” |
| 3. | / atampenda / | “he will like him” |
| 4. | / nitakupenda / | “I will like you” |
| 5. | / nitampenda / | “I will like him” |
| 6. | / utanipenda / | “you will like me” |
| 7. | / utampenda / | “you will like him” |
| 8. | / tutampenda / | “we will like him” |
| 9. | / watampenda / | “they will like him” |
| 10. | / atakusumbua / | “he will annoy you” |
| 11. | / unamsumbua / | “you are annoying him” |
| 12. | / atanipiga / | “he will beat me” |
| 13. | / atakupiga / | “he will beat you” |
| 14. | / atampiga / | “he will beat him” |
| 15. | / ananisumbua / | “he is annoying me” |
| 16. | / anakusumbua / | “he is annoying you”. |
| 17. | / anakulipa / | “he is paying you” |
| 18. | / amekulipa / | “he has paid you” |
| 19. | / atakulipa / | “he will pay you” |
| 20. | / alinipenda / | “he liked me” |
| 21. | / alikupenda / | “he liked you” |
| 22. | / alinilipa / | “he paid me” |
| 23. | / wamekulipa / | “they have paid you” |
| 24. | / tutakulipa / | “we will pay you” |
| 25. | / umenipenda / | “you have liked me” |
| 26. | / tulimpenda / | “we liked him” |

(Do not be misled by word order in the gloss, or think that each word in the gloss must correspond to a morpheme in the item—for example the gloss “is/areing” corresponds to one morpheme as does the gloss “has -ed”. Also do not be misled by changes in the gloss which are peculiarities of English—for example the two glosses “has” and “have” correspond to the same morpheme in the corpus).

A: List all the morphemes and give their approximate meanings:

B: Explain what problem occurs with the phoneme string / ni /, how can you account for its apparent change of gloss (hint—perhaps they are different morphemes, or perhaps they have different positional meanings):

4–4 Tepehua (Amerind, Mexico) (Adapted from Gleason, *Workbook in Descriptive Linguistics*, 1Ed. ©1955:27 Wadsworth, a part of Cengage Learning, Inc. Reproduced by permission. www.cengage.com/permissions)

- | | | | |
|---------------------|----------|--------------------------------|----------------|
| 1. / laqatam / | “one” | 9. / laqakaawtʰutu / | “thirteen” |
| 2. / laqatʰuy / | “two” | 10. / laqapʰuʂam / | “twenty” |
| 3. / laqatʰutu / | “three” | 11. / laqapʰuʂamtam / | “twenty one” |
| 4. / laqatʰaatʰii / | “four” | 12. / laqapʰuʂamkaaw / | “thirty” |
| 5. / laqakiis / | “five” | 13. / laqapʰuʂamkaawtʰuy / | “thirty two” |
| 6. / laqakaaw / | “ten” | 14. / laqapʰuʂamkaawnahaač / | “thirty nine” |
| 7. / laqakaawtam / | “eleven” | 15. / laqatʰaatʰiikiispʰuʂam / | “four hundred” |
| 8. / laqakaawtʰuy / | “twelve” | 16. / laqanahaačkiispʰuʂam / | “nine hundred” |

Note: double vowel graphs = long vowel; / pʰ / and / tʰ / =glottalized stops

List all the morphemes and give their approximate meanings (glosses):

What problem is present in assigning all phoneme sequences to morphemes?

Does the placement of a morpheme in a word (item) affect its significance to the meaning of the word? (Explain)

What items would you predict for the following glosses:

- “nine” / _____ /
- “one hundred and two” / _____ /
- “forty” / _____ /

4–5 Turkish (Middle East) (Adapted from Gleason, *Workbook in Descriptive Linguistics*, 1Ed. ©1955:35 Wadsworth, a part of Cengage Learning, Inc. Reproduced by permission. www.cengage.com/permissions)

- | | | | |
|---------------|-----------------|-----------------|------------------------|
| 1. / ziller / | “bells” | 9. / zilimiz / | “our bell |
| 2. / zilli / | “with the bell” | 10. / zilin / | “your (singular) bell” |
| 3. / zilden / | “from the bell” | 11. / ziliniz / | “your (plural) bell” |

4. / zillerden / “from the bells” 12. / zillerinizden/ “from your (plural) bells”
 5. / zilim / “my bell” 13. / zilinde / “in your (singular) bell”
 6. / zillerim “my bells” 14. / zilime / “to my bell”
 7. / zilimden / “from my bell” 15. / zillerimize / “to our bells”
 8. / zillerimden “from my bells” 16. / zilimizli “with our bell”

List all the morphemes and give their approximate meaning:

4–6 Luganda (Uganda, East Africa) (Adapted from Gleason, *Workbook in Descriptive Linguistics*, ©1955:24 Wadsworth, a part of Cengage Learning, Inc. Reproduced by permission. www.cengage.com/permissions)

- | | |
|-------------------------|--------------------------|
| 1. / omukazi / “woman” | 6. / abakazi / “women” |
| 2. / omusawo / “doctor” | 7. / abasawo / “doctors” |
| 3. / omusika / “heir” | 8. / abasika / “heirs” |
| 4. / omuwala / “girl” | 9. / abawala / “girls” |
| 5. / omulenzi / “boy” | 10. / abalenzi / “boys” |

List all the morphemes and give their approximate meanings.

What problem is present in determining the gloss in English for a morpheme such as / sawo /?

4–7 Klingon (artificial, science fiction, from materials in Okrand)

- | | |
|---|---|
| 1. / ɖoʂɖaɟ / “his target” | 8. / ɖuʂme ⁱ ɖaɟ / “his (torpedo) tubes” |
| 2. / yuq ^h ɖaɟ / “his planet” | 9. / ninliɟ / “your [singular] fuel” |
| 3. / ɖeʂɖaɟ / “his arm” | 10. / yaʂme ^{li} ? / “your [sing.] officers” |
| 4. / yuq ^h wiɟ / “my plane” ^t | 11. / mebwɪ? / “my guest” |
| 5. / ɖoʂme ⁱ wiɟ / “my targets” | 12. / yaʂli? / “your [singular] officer” |
| 6. / ɖujwiɟ / “my (space) ship” | 13. / puqli? / “your [singular] child” |
| 7. / ɖuʂwiɟ / “my (torpedo) tube” | 14. / puqme ⁱ wi? / “my children” |

Note: ɖ and ʂ are articulations made with the tongue tip moved further back toward the palate

List all the morphemes and give their approximate meanings.

What assumption(s) did you have to make about embedded meanings in order to determine all the morphemes?

4–8 Hebrew (Middle East, drawings by Roger Nevarez) (Adapted from Gleason, *Workbook in Descriptive Linguistics*, 1Ed. ©1955:9 Wadsworth, a part of Cengage Learning, Inc. Reproduced by permission. www.cengage.com/permissions).

These four drawings represent four situations. Below them are labels giving verbs that might be used to describe them in English and Hebrew. What is the difference in how these are used—explain what a translator must know in order to translate these Hebrew verbs into English.



English: “comes” Hebrew: / bá: / English: “comes” Hebrew: / ya:sá: /



English: “goes” Hebrew: / ya:sá: English: “goes” Hebrew: / bá: /

In translating Hebrew / bá: / and / ya:sá: / into English how does one determine whether to use “come” or “go”?

4–9 Look up one of the following words in a reasonably complete dictionary (must be print, not online) and answer the following questions about one:

nut star empty objective solid camp soil

Describe how your dictionary handles the word (is there just one entry or more than one, are different “meanings” presented under one entry, and how are they distinguished, are some meanings identified as more important than others...):

For this word, based upon your own usage, do you believe that the “first” meaning listed is the denotative meaning? Explain.

For this word, and based on your own usage, are all the meanings included under one entry connotative meanings associated with one denotative meaning or should some have been given a separate entry? Explain.

4–10 Construct a semantic “map” of one of the following domains (circle your choice):

club meeting

bar (place to drink)

dining room (in residence)

garage (in residence)

What morphemes are **typically** associated with this domain?

Categorize these morphemes according to their semantic function (nouns, adjectives) and relationships (co-occurrence, sequence, taxonomies):

Exercises for Chapter Five

5–1 Stem and affix identification in Klingon (artificial). Refer back to the corpus for exercise 4–7 (page 150) on Klingon morpheme identification. The corpus contains many stems and affixes, however each item consists of one stem and several affixes.

Using at least two of the criteria set out in chapter 5, identify the stem morphemes and list these with their approximate meaning (gloss):

Explain how you made these decisions (which criteria did you use—these must have been used for all).

List the affixes with their approximate meaning (gloss) and state which kind of affix each is:

5–2 Affix Orders (Turkish)

Given that /zil/ “bell” is the stem in the following corpus, identify the affix morphemes (they are all suffixes) and then group them into members of affix orders. Provide an appropriate semantic label for each order.

- | | | | |
|------------------|------------------|---------------------|---------------------------|
| 1. /ziller/ | “bells” | 10. /zilimiz/ | “our bell” |
| 2. /zilli/ | “with the bell” | 11. /zilin/ | “your (sing.) bell” |
| 3. /zilden/ | “from the bell” | 12. /ziliniz/ | “your (plur.) bell” |
| 4. /zillerden/ | “from the bells” | 13. /zillerinizden/ | “from your (plur.) bells” |
| 5. /zilim/ | “my bell” | 14. zilinde | “in your (sing.) bell” |
| 6. /zillerim/ | “my bells” | 15. /zilime/ | “to my bell” |
| 7. /zillerimiz/ | “our bells” | 16. /zillerimize/ | “to our bells” |
| 8. /zillerimden/ | “from my bells” | 17. /zilimizli/ | “with our bell” |
| 9. /zilimden/ | “from my bell” | | |

There are four affix orders, give the affixes that occur in each order, its gloss and the semantic label for each order (may not be much different):

5–3 Allomorphic variation in plural morphemes (English). Using yourself as an informant, determine the phonemic form of the plural suffix for each of the following noun stems (the plural suffix for the first noun is given as an example).

<u>word</u>	phonemic form for <u>stem*</u>	phonemic form for <u>plural affix</u>	<u>word</u>	phonemic form <u>for stem*</u>	phonemic <u>for plural affix</u>
“cat”	/kæt/	/-s /	“lathe”		
“dog”	/ -g /	/-z /	“bar”		
“lap”			“catch”		
“lab”			“judge”		
“eye”			“ham”		
“zoo”			“moon”		
“cliff”			“hall”		
“love”			“gash”		
“truck”			“garage”		
“myth”			“buzz”		
“feud”			“song”		
“axe”					

*Only need to give the final phoneme in the stem

Determine what the pattern is for the variation in these plural suffixes:

1) state what kind of allomorphic variation is occurring; 2) state the rule for allomorphic variation of these English plural suffixes (i.e. “the plural suffix takes the form / ____ / in X environment; it takes the form /

___ / in Y environment; and so on (for as many different allomorphs as there are in the corpus).

5–4 Turkish (Middle East) (Adapted from Gleason, *Workbook in Descriptive Linguistics*, 1Ed. ©1955:35 Wadsworth, a part of Cengage Learning, Inc. Reproduced by permission. www.cengage.com/permissions)

Given: /adam/ “man”, /baş/ “head”, /el/ “hand”, /ev/ “house”, /kol/ “arm”, /zil/ “bell”, /ses/ “voice”, and /diş/ “tooth”.

- | | |
|-------------------------------|----------------------------------|
| 1. /adama/ “to the man” | 11. /baştan/ “from the head” |
| 2. /başa/ “to the head” | 12. /adamlar/ “men” |
| 3. /ele/ “to the hand” | 13. /başlar/ “heads” |
| 4. /eve/ “to the house” | 14. /eller/ ¹ “hands” |
| 5. /kola/ “to the arm” | 15. /evler/ “houses” |
| 6. /zile/ “to the bell” | 16. /kollar/ “arms” |
| 7. /koldan/ “from the arm” | 17. /ziller/ “bells” |
| 8. /zilden “from the bell” | 18. /kollardan/ “from the arms” |
| 9. /sesten/ “from the voice” | 19. /zillerden/ “from the bells” |
| 10. /dişten/ “from the tooth” | |

¹Note: double consonant graphs indicate two identical consonants in succession

Determine each group of affix allomorphs and explain why they can each be considered allomorphs of the same morpheme:

State the rules for the allomorphic variation within each morpheme.

What kind of allomorphic variation is this?

Given /kuş/ “bird”, what form would you predict for the gloss “from the birds”?

5–5 Comparison of Morphologically Distributed Allomorphs for Dinka (East Africa, Sudan) and English

Dinka

singular form

/ yot / “hut”
 / bit / “spear”
 / agook / “monkey”
 / tuoŋ / “egg”
 / yič / “ear”
 / gol / “clan”
 / met / “child”

plural form

/ yoot / “huts”
 / biit / “spears”
 / agok / “monkeys”
 / toŋ / “eggs”
 / yit / “ears”
 / gal / “clans”
 / miit / “children”

Englishsingular form

/alumnəs/

/daⁱ/

/aks/

/mæn/

“alumnus”

“die”

“ox”

“man”

plural form/alumnaⁱ//daⁱs/

/aksən/

/men/

“alumni”

“dice”

“oxen”

“men”

Compare these two systems. In general, how are most of the Dinka and English plural forms different from each other? Which English form is similar to those of the Dinka, and why?

5–6 Syntax, Word order (English)

With a frame such as:

“John bought _____ house”

many modifiers (adjectives) may be used by English speakers. Examine the list below of a sample of these and do the following:

A. Group these words into six modifier sub-classes according to their ability to substitute for each other (i.e. the members of a sub-class cannot be used together, only one can be used in the appropriate frame):

B. What is the order of occurrence of these six sub-classes. (This is very much like figuring out affix orders. That sub-class that can occur next to the noun it is modifying is sub-class #1, that which occurs next to #1 is #2, and so on.)

a (an)	large	small	attractive	modern	stone
bad	my	the	blue	little	white
brown	new	wood(en)	brick	nice	this
good	old	your	his		

5–7 Write a generative grammar that will yield the following English sentences, *and only these*.

1. “The child loves the toy”
2. “The child will love the toy”
3. “The toy is loved by the child”
4. “The toy will be loved by the child”

Start with the mental entity S and then list the necessary replacement and transformational rules that will generate these sentences, and only these sentences. You must include a lexicon but you do not have to include any phonological rules.

SUGGESTED SOLUTIONS TO EXERCISES

Chapter Two

2–4 Examples of four household/dorm room objects

<u>letters</u>	<u>manner</u>	<u>place of articulation for consonants</u>	<u>place for vowels</u>
r	resonant	palatal	mid central
u	V		
g	stop	velar	
“desk”			
d	stop	alveolar	mid front
e	V		
s	fricative	alveolar	
k	stop	velar	
“light”			
l	resonant	alveolar	low central then high front
i	V		
ght	stop	alveolar	
“notes”			
n	resonant	alveolar	mid back (then high back)
o	V		
t	stop	alveolar	
es	fricative	alveolar	

2–5 Examples selected from both lists.

“lamp”	l	a	m	p
manner	resonant	V	resonant	stop
voicing	vd		vd	vl
aspiration	unasp		unasp	unasp
place	alveolar		bilabial	bilabial
“glove”	g	l	o	ve
manner	stop	resonant	V	fricative
voicing	vd	vd		vd
aspiration	unasp	unasp		unasp
place	velar	alveolar		labiodental
“scrub”	s	c	r	u
manner	fricative	stop	resonant	V
voicing	vl	vl	vd	vd
aspiration	unasp	unasp	unasp	unasp
place	alveolar	velar	palatal	bilabial

“mail”	m	ai	l	
manner	<i>resonant</i>	<i>V</i>	<i>resonant</i>	
voicing	<i>vd</i>		<i>vd</i>	
aspiration	<i>unasp</i>		<i>unasp</i>	
place	<i>bilabial</i>		<i>alveolar</i>	
“judge”	j	u	d	je
manner	<i>affricated stop</i>	<i>V</i>	<i>stop</i>	<i>affricated stop</i>
voicing	<i>vd</i>		<i>vd</i>	<i>vd</i>
aspiration	<i>unasp</i>		<i>unasp</i>	<i>unasp</i>
place	<i>alveopalatal</i>		<i>alveolar</i>	<i>alveopalatal</i>

2–6 Using the same words as in 2 – 5, this is a description of the vowel characteristics

“lamp”		l	a	m	p
complex/simple		<i>C</i>	<i>simple</i>	<i>C</i>	<i>C</i>
height			<i>low</i>		
position			<i>front</i>		
rounding			<i>unro</i>		
nasalization			<i>nas</i>		
length			<i>(not applicable)</i>		
“glove”	g	l	o	ve	
complex/simple	<i>C</i>	<i>C</i>	<i>simple</i>	<i>C</i>	
height			<i>mid</i>		
position			<i>central</i>		
rounding			<i>unro</i>		
nasalization			<i>nonnas</i>		
length			<i>(not applicable)</i>		
“scrub”	s	c	r	u	b
complex/simple	<i>C</i>	<i>C</i>	<i>C</i>	<i>simple</i>	<i>C</i>
height				<i>mid</i>	
position				<i>central</i>	
rounding				<i>unro</i>	
nasalization				<i>nonnas</i>	
length				<i>(not applicable)</i>	
“mail”	m		ai		l
complex/simple	<i>C</i>		<i>complex (forward glide)</i>		
height			<i>mid to high</i>		
position			<i>front</i>		
rounding			<i>unro</i>		
nasalization			<i>nonnas</i>		
length			<i>(not applicable)</i>		
“judge”	j	u	d	je	
complex/simple	<i>C</i>	<i>simple</i>	<i>C</i>	<i>C</i>	

height	<i>mid</i>
position	<i>central</i>
rounding	<i>unro</i>
nasalization	<i>nonnas</i>
length	<i>(not applicable)</i>

2–7 Significant differences (suggested answers in *italics*)

stop—fricative: *for fricative airstream is not interrupted as it is for a stop*
 alveolar—alveo-palatal: *the tip of the tongue is the moveable articulator for alveolar but in alveopalatal it is the front of the tongue*; median resonant—lateral resonant: *the tongue is touching the top of the oral chamber for a lateral resonant but not for a median resonant*; aspirated stop—affricate: *aspirated stop has airstream released with an extra “puff,” affricated stop has airstream released into a narrow (fricative-like) opening*; front V—back V: *front part of tongue is highest part in front vowel while back part of tongue is highest for back vowel*; forward glide—back glide: *tongue moves forward in oral chamber for forward glide and backwards for a back glide*; resonant—vowel: *airstream is redirected to a greater or lesser degree for a resonant but for a vowel the airstream is shaped by the dimensions of the oral chamber*; bilabial—labiodental *bilabial uses both lips while the labiodentals uses the lower lip and the upper front teeth*; velar stop—glottal stop: *velar stop involves the tongue back while the glottal stop uses both of the laryngeal folds to interrupt the airstream*; voiced C—voiceless C: *voiced C has vibration of the vocal cords as a co-articulation, voiceless C does not*; bilabial stop—bilabial nasal resonant: *uvula is raised to close off nasal chamber for a bilabial stop but is lowered for a bilabial nasal resonant*; high V—low V: *tongue is relatively high in the oral chamber for a high V and low for a low V* complex V (glide)—simple V *tongue configuration changes during a complex V but remains the same during a simple V*; rounded V unrounded V: *during a rounded V the lips are tightened (puckered) but the lips are relaxed during an unrounded V*; nasalized V—nonnasalized V: *the uvula is lowered during a rounded V so that the airstream goes though both the oral and nasal chambers but the uvula is raised for an oral V so that the airstream only goes though the oral chamber.*

2–8 A Phonetic notation – consonants

Words selected from each list.

<u>Word</u>	<u>phonetic spelling</u>	<u>Word</u>	<u>phonetic spelling</u>
“bank”	[bVŋk]	“once”	[wVns]
“ghost”	[gVst]	“three”	[ΘjV]
“schedule”	[skVdʒVl]	“enough”	[VnVf]

2–8 B Phonetic notation – vowels (same words)

<u>Word</u>	<u>phonetic spelling</u>	<u>Word</u>	<u>phonetic spelling</u>
“bank”	[CæCC]	“once”	[CəCC]
“ghost”	[Co ^w CC]	“three”	[CCi ⁱ]
“schedule”	[CCeCCəC]	“enough”	[ɪ ⁱ CəC]

2–9 Polysyllabic words

<u>Word</u>	<u>phonetic spelling with prosodic features (stress)</u>
“linguistics”	[li ⁿ ŋgwístiks]
“phonetics”	[fō ^w nétiks]

2–10 Phonetic transcription of speech segment

[æsk náʔ wət yu^wʔ kəntʃiⁱ kæn du^w fo^wʔ yú^w

æsk wət yú^w kæn du^w fo^wʔ yu^wʔ kəntʃiⁱ]

“Ask not what your country can do for you,
ask what you can do for your country”

Chapter Three

3–1 Suggested examples of minimal pairs. See if you can come up with others. Remember these are based on the author’s dialect of American English, yours may be different—just don’t rely on spelling!

Speech sound		position in word	
<u>pair</u>	<u>initial</u>	<u>medial</u>	<u>final</u>
[p] – [b]	none	“rapid” [jæpid]	“lip” [lip]
		“rabid” [jæbid]	“lib” [lib]
[f] – [v]	“fine” [f ā ⁱ n]	“define” [di ⁱ f ā ⁱ n]	“half” [hæf]
	“vine” [v ā ⁱ n]	“devine” [di ⁱ v ā ⁱ n]	“have” [hæv]
[f] – [θ]	“first” [fɪrst]	“offer” [əfeɪ]	“deaf” [def]
	“thirst” [θɪrst]	“author” [əθəɪ]	“death” [deθ]
[s] – [ʃ]	“sip” [sɪp]	“mast” [mæst]	“bass” [bæs]
	“ship” [ʃɪp]	“mashed” [mæʃt]	“bash” [bæʃ]
[s] – [z]	“sip” [sɪp]	“re-sign” [ri ⁱ ˈsā ⁱ n]	“mace” [meɪs]
	“zip” [zɪp]	“resign” [ri ⁱ ˈzā ⁱ n]	“maize” [meɪz]
[k] – [g]	“clue” [klu ^w]	“bicker” [bɪkeɪ]	“duck” [dæk]
	“glue” [glu ^w]	“bigger” [bɪgeɪ]	“dug” [dæg]
[n] – [ŋ]	none	“sinner” [sɪneɪ]	“kin” [kɪn]
		“singer” [sɪŋeɪ]	“king” [kɪŋ]
[i] – [æ]	“it” [ɪt]	“knit” [nɪt]	none
	“at” [æt]	“gnat” [næt]	
[e ⁱ] – [a ⁱ]	“ache” [e ⁱ k]	“Dave” [deɪv]	“day” [deɪ]
	“Ike” [a ⁱ k]	“dive” [daɪv]	“die” [daɪ]

3–2 Zoque vowels

UNROUNDED

		Location of Tongue Height					
		<u>Front</u>		<u>Central</u>		<u>Back</u>	
		nasal	non-nasal	nasal	non-nasal	nasal	non-nasal
Tongue Height	High		i				
	Lower-high						
	Mid		e		ə		
	Lower-mid						
	Low				a		

ROUNDED

		Location of Tongue Height					
		<u>Front</u>		<u>Central</u>		<u>Back</u>	
		nasal	non-nasal	nasal	non-nasal	nasal	non-nasal
Tongue Height	High						
	Lower-high						u
	Mid						o
	Lower-mid						
	Low						

3–3 Chamorro consonants

	bi- labial	labio- dental	dental	alveolar	alv- palatal	palatal	velar	uvular	glottal
	vl vd	vl vd	vl vd	vl vd	vl vd	vl vd	vl vd	vl vd	vl
simple stops	p b			t d			k g		ʔ
aspirated stops									
affricated stops									
fricatives		f		s					h
nasal									
<u>resonants</u>	m			n			ŋ		
lateral resonants				l					
median resonants									

3–4 Spanish distributional analysis for [β] and [b]

Listing preceding and following environments for each sound:

	[β]			[b]	
a		a	#		o
o		o			a ⁱ
i		ε			i
e					ε
ĩ					

Are these two sounds in complementary distribution? *yes*; Describe the distribution: *no identities in preceding environment, [b] preceded by silence, [β] preceded by vowels and a flap*; What is their phonemic status? *they are probably allophones of the same phoneme*

3–5 Chiricahua Apache distributional analysis for [ɣ̤] and [ɣ]

	[ɣ̤]			[ɣ]	
#		e:	#		o
s		e	ɬ		a
ɬ		i			
t		i:			

Are these two sounds in complementary distribution? *Yes*; Describe the distribution: *no identities in following environment, [ɣ̤] followed by front*

vowels while [x] is followed by central and back vowels; What is their phonemic status? *They are probably allophones of the same phoneme*

3-6 Kiswahili Distributional analysis of [o] and [ɔ]

	[o]			[ɔ]	
d		t	g		m
t		#	b		n
ʃ		g	u		ñ
g		k	ñ		
u		ʃ	k ^h		
k					
m					

Are these two sounds in complementary distribution? *yes*; Describe the distribution: *no identities in following environments, [ɔ] occurs followed by nasal resonants, [o] occurs followed by non-nasal resonants*; What is their phonemic status? *They are probably allophones of the same phoneme*

3-7 Chamorro Distributional Analysis of [n] and [m]

	[n]			[m]	
#		u	#		a
u		i	a		ɪ
a		o	o		#
		#	u		o
		a			u
					i

Are these two sounds in complementary distribution? *No*; Describe the distribution: *Identities in both preceding and following environments (Also, presence of minimal pair items # 4 and # 7)*; What is their phonemic status? *Must be members of different phonemes*

3-8 Manx Distributional analysis of [k] and [g]

	[k]			[g]	
a		#	a		#
#		ɛ	#		ɛ
ɛ		e	ɛ		æ
ɪ		r	ɪ		r
		l			o
		æ			

Are these two sounds in complementary distribution? *No*

Describe the distribution *Identities in both preceding and following environments.*

What is their phonemic status? *Must be members of different phonemes.*

3–9 Zoque allophonic consistency

	[p]			[b]	
x̣		u	n		a
#		i	ŋ		y
t		e	m		
ə		a			
		o			
		y			
		ə			
		ḳ			

[p] and [b] in comp. distrib. therefore are allophones of same phoneme; Prediction: voiced and voiceless stops are in allophonic relationships. Vd stop allophone occurs when preced. environment is a nasal resonant, otherwise vl stop allophone occurs.

Testing prediction using [tʲ] and [dʲ]

	[tʲ]			[dʲ]	
ə		u	ñ		u
#		ə	m		o
			ŋ		a

[tʲ] and [dʲ] also in comp. distribubution; In general, voicing does not appear to be a phonemic boundary for Zoque stops.

3–10 Determination and application of phoneme graphs for Zoque

Basis for selection of phoneme graphs: *I will select the voiceless allophone phonetic graphs in each phoneme to use as the graphs for that phoneme. This is based on the fact that the voiceless allophone occurs in more kinds of environments*

/ p / [p, b], / t / [t, d], / tʲ / [tʲ, dʲ], / ts̄ / [ts̄, d̄z], / č̄ / [č̄, ʝ], / k / [k, g]

item rewrites

#3= /kaŋ/, #19= /n̄tsima/, #29= /ŋkəʔ/, #52 = /čehčaxu/, #4= /kaʔnči /

#22= /ñčeht̄su/, #34= /petpa/, #12= /mpata/, #22= /nətʲuxu/, #45= /tiŋtiŋ

Chapter Four

4–1 Suggested morpheme identification for Classical Nahuatl

*Examine the whole corpus for co-occurrences of item phoneme strings with gloss elements: “stone” occurs in the glosses of #1 and #7 and the phoneme string /tew/ co-occurs as well in each item. “stone” does not show up in any other gloss, and /tew/ does not show up in any other item. This results in the likely identification of /tew/ as a morpheme with the gloss “stone”. Use the same method to identify the likely identifications of: /petʰ/ “mat,” /taʔ/ “father,” /kwik/ “song,” /tʰew/ “fire,” /aamaw/ “book;” Two phoneme strings are left unidentified—/no/ and /ii/ and these can fairly readily be seen as co-occurring with the respective glosses “my” (items #1–6) and “his” (items #7–12); There are therefore eight likely morphemes: /petʰ/ “mat,” /taʔ/ “father,” /kwik/ “song,” /tʰew/ “fire,” /aamaw/ “book,” /no/ “my,” /ii/ “his;” Keep in mind that **all** phoneme strings must be linked to gloss elements, you cannot have any “left over.”*

4–2 Suggested morpheme identification for Delaware

The gloss “little” for items #1–6 co-occurs with the phoneme string /tət/ and is a likely morpheme identification. You should also be able to see that “father” co-occurs with /uux/, “foot” with /siit/, “mouth” with /tuun/, and “daughter” with /taan/. This leaves three phoneme strings (of one phoneme length) left. /n/ co-occurs with “my,” /k/ with “your,” and /w/ with “his.” The likely morphemes you can identify from this corpus are therefore: /tət/ “little,” /uux/ “father,” /siit/ “foot,” /taan/ “daughter,” /n/ “my,” /k/ “your,” /w/ “his.”

4–3 Suggested morpheme identification for Kiswahili

This problem presents some difficulties due to the English glosses. There are different tenses in the English glosses (e.g. “ed” for [past] and “is/are...ing” for [present]) but note that these are probably not the same in Kiswahili. Consequently, unlike the previous problems, it might be better to start by finding repeated phoneme strings rather than gloss elements. Notice the common string /penda/ in items #1–9. The common gloss element is “will like” but in items #20, 21, 25, 26 the “will” (or [future tense]) is not present, rather “ed (or [past tense])”. Therefore the best gloss for /penda/ would be “like (having no tense)”. Using the same reasoning you should be able to get /sumbua/ “annoy (having no tense)”, /piga/ “beat (having no tense)”, and /lipa/ “pay (having no tense).”

Now for other phoneme strings that recur: There are a large number of items that begin with /a/ (items #1–3, 10, 12–22) and these co-occur

with the gloss element “he”, so you can also identify a likely morpheme /a/ with the gloss “he”. Similarly, /ni/ co-occurs with “I” in items #4 and 5 (do not be misled by the fact that there is also a /ta/ string in these items as well—/ta/ also occurs in other items, e.g. #3, but not with “I” so it must be a different morpheme). You should be able to continue and identify the likely morpheme glosses for the strings /u/, /tu/, and /wa/.

Now consider the string /ta/ occurring in items # 1—10, 12—14, 19 and 24. Each gloss has “will” as an element so this string could likely represent a [future] tense, or “will”. The string /na/ co-occurs with “is/are ...ing” or a [present] tense, /i/ with “ed” or a [past] tense, and /me/ with “have/has...ed” or a [past perfect] tense. (Again, do not be put off by the fact that one string in Kiswahili is a morpheme with several strings in the gloss—this is the challenge of using one language to map another.) There are still two phoneme strings not accounted for, /-m-/ (items # 3, 5-7, 11, 14, 26) and /-ku-/ (items 2, 4, 10, 13, 16—19, 21, 23, 24) and these, respectively, can be tied to the glosses “him” and “you”. A problem arises in that there are two “you” glosses. This a problem with the fact that English used one phoneme string for both singular and plural second persons, and for both a subject and an object. Consequently, you should indicate the gloss for /ku/ as “you (object)”.

B. The Kiswahili phoneme string /ni/ occurs with two glosses—“I” and “me”. Both are first person, but one is for a subject and the other for an object. One way to treat these is as homophones—two separate morphemes that happen to have the same phoneme string, a fairly common feature of all languages, and they would be listed twice in the list of morphemes. Another way is to treat them as the same morpheme whose meaning is based upon its position in the word. /ni/ = “I” if it occurs before the tense morpheme and as “me” if after. English does somewhat the same with its “you” /yu/ morpheme.

4-4 Tepehua morpheme identification

The problem appears from the glosses to be a straightforward correlation of phoneme strings in the items with “numbers” in the glosses, however you will have to consider some embedded numeric elements in the glosses. Let’s start with item #1 and its gloss of “one.” The initial phoneme string of /ʌaqa/, however, recurs for all items in the corpus and so, for now, we must set it aside as related to a particular number gloss such as “one”. This leaves /tam/ and you must see if shows up with other “one” glosses. #11 “twenty one” is one such occurrence and indicates that /tam/ is a morpheme with the gloss of “one.” However item #7 also has the /tam/ string but does not appear to have “one” in the gloss, and this would damage using it with the gloss of “one.” Consider that while

the English gloss for #7 is “eleven”, the embedded meaning is something like “ten plus one” thus providing a “one” gloss. (You should not push this approach too far, but here it seems reasonable.) This same approach can be used with the strings /^ʔuy/ and /^ʔutu/, for “two” and “three” respectively, using items #2, 8, 13, and #3, and 9. In addition, this approach should yield identifying /kaaw/ with “ten”. Given the model of items #1 – 5, you could tentatively identify /^ʔaatii /, /kiis /, and /p^ʔu^ʂam/ as morphemes with the glosses “four”, “five”, and “twenty” respectively. However it would be better if you could find other occurrences of each with the same gloss. Items #15 and 16 provide this, but note that there is a difference in meaning. Consider again items #7 and 8 where a number following /kaaw / “ten” is added to it. In items #15 and 16 /^ʔaatii/ and /kiis/ are multiplied. “four hundred” can be understood as being “four” [times] “five” [times] “twenty”. Thus item #15 provides cooccurrence for /^ʔaatii/, /kiis /, and /p^ʔu^ʂam /. You should also be able to identify the morpheme for “nine” from items #14 and 16.

Rather than treat the two possible meanings of [number plus] and [number times] as homophones, it seems better to treat the present of [addition] or [multiplication] as positional extended meanings of numbers: it appears that if a number morpheme precedes /p^ʔu^ʂam / it is multiplied (e.g. [one hundred] is /kiis/ “five” preceding /p^ʔu^ʂam/ “twenty”). If a number follows /p^ʔu^ʂam/ or /kaaw / it is added. Therefore for your listing of the morphemes you should add the parenthetical statement following the number gloss of something like “[number] (added or multiplied based on position in the word)”.

/laqa / presents a problem since it is never absent. You would need its absence to see what happens to the gloss. Some guess that it means [number], but this cannot be determined from this corpus. You could include it in your morpheme list but with only a parenthetical statement such as “(undetermined reference, occurs with all items)”

Now you should be able to predict the items that would exist for the glosses “nine”, “one hundred and two”, and “forty” (we will omit /laqa/): /^ʔna^ʂaa^ʂ/ “nine,” /kiisp^ʔu^ʂam^ʔuy/ “one hundred and two” (some place the number /tam/ in front of /kiis/), /^ʔuy^ʔp^ʔu^ʂam/ “forty” (some have suggested /p^ʔu^ʂamkaawkaaw/ or even /p^ʔu^ʂamp^ʔu^ʂam/, but I do not think there is enough evidence for these forms).

4–5 Identification of Turkish morphemes

Somewhat like the previous problem (in which a phoneme string was always present but with no direct association with a gloss element), here the constant phoneme string, /zil / is always present in cooccurrence with the gloss element “bell.” Although without its absence and a corresponding absence of “bell” in the gloss, it is cannot be a definite

identification. You could list /zıl/ in your morpheme list but you would have to add a parenthetical qualification—"bell (but not sure since never absent)". You should be able to see that /ler/ co-occurs with "-s (plural)". Now, either starting with gloss elements or phoneme strings, you should be able to see that /den/ co-occurs with "from". (NOTE--the presence of "the" in the gloss is undoubtedly only due to a convention of English to require an article since there is no phoneme string that corresponds to it—compare items #1 and 4. You might want to list /den/ with the gloss "from the" but be aware that the article only occurs in a singular gloss, e.g. item #2, and also not when there is a possessive pronoun, e.g. #8.) From items #2 and 16 you can associate /li/ with "with", from #14 and 15 /e/ with "to", and, although there is only one occurrence, #13, /de/, with "in".

The possessive pronouns will require that you use your knowledge of embedded meanings in the English gloss. One approach might be to list /im/ with the gloss "my" and /imiz/ with "our". However the presence of the /im/ string with both "my" and "our" glosses, and especially the presence of /iz/ with both "our" and "your (plural)", should make you suspicious. The presence of the same phoneme string in different morphemes is not unusual, for example the phoneme strings /de/ and /den/ above share a /de/ sequence. But for these there is no apparent semantic linkage. But for /im/, /imiz/, and /iniz/ it should occur to you that the /im/ and /in/ sequences refer to the pronoun person number (first person and second person, respectively) and the /iz/ sequence pluralizes the pronoun—added after /im/ it changes "my" to "our" and added to /in/ it changes "your (singular)" to "your (plural)".

It is interesting that there are two plural morphemes in Turkish, one for the noun and the other for the pronoun. Since English has two corresponding morphemes you will have to use parenthetical explanations in your morpheme glosses: /im/ "(first person possession for pronoun)," /in/ "(second person possession for pronoun)," /iz/ "(pluralizes possessive pronoun)."

4-6 Luganda morpheme identification

This problem seems straightforward. /omu/ appears with the singular form of a noun and /aba/ with the plural form. Each noun is then apparently revealed when these prefixes are removed—e.g. /kazi/ means "woman". However, as in the previous problem these noun glosses would not be really accurate. In English there is a [plural] morpheme but no singular morphemes (there are a few exceptions, e.g. *alumnus* and *alumni*), and so those morphemes designating [entities], or "nouns," such as "woman" are also embedded with [singular]. The English plural morpheme thus removes the [singular] reference and replaces it with

[plural]. But Luganda has a singular and a plural morpheme, and this means that “nouns” such as /kazi/ properly have no number by themselves. Consequently representing /kazi/ with the gloss “woman” is not correct since “woman” is also singular but /kazi/ is not. Consequently, you must add a parenthetical statement to the gloss—e.g.

/kazi/ “(woman, but neither singular nor plural)”

4–7 Klingon morpheme identification

You may begin by looking for repeated gloss elements or item phoneme strings. Let’s use gloss elements and note the repeated presence of “his” in items #1–3 and 8. Co-occurring in the items is the string /ɖaɰ/. However /ʝ/ also occurs in items #4–7 and 9 with no “his” in the gloss. In particular compare items #2 and 4. The only difference in the gloss is the change from “his” to “my” and in the item the only difference is the change from /ɖa/ to wi/. /ɖa/ can be identified as a morpheme with the gloss “his”. Items #4 and 11 yield /wi/ as “my”, and items #9, 10, 13 and 13 yield /i/ as “your (singular)”. Items # 5, 8, 10, and 14 have plural for the noun and the string /me/ only occurs for these items, consequently /me/ may be listed as a morpheme with a gloss of “(plural)” (you might want to just give the gloss as “-s” but the presence of “-ren” in item #14 prevents this). With the noun plural and the possessive pronouns identified, and ignoring the final phoneme in each item (either /ʝ/ or /ʔ/, identifying the nouns should be easy, for example items #1 and 5 give /ɖos/ as “target”. However three probably nouns only occur once and so for those (“arm”, “fuel”, and “guest”) you will have to add a parenthetical statement such as “...(identification not definite since it only occurs once in the corpus)”.

Now we can turn to the final phoneme in each item, either /ʝ/ or a /ʔ/. You will need to look at embedded meanings in the glosses. As a hint look at the gloss for item #4 and compare its reference to that of item #11. What is being possessed? English makes no distinction in its possessive pronouns as to what is being possessed, but many languages do. In English I can say “my book” or “my child,” but are these the same kinds of “possessing?” I can sell my book but not my child. One of these two phonemes refers to one kind of possession, of the nature of what is “possessed” and you should be able to figure this distinction out. You will have to use parenthetical descriptions in your glosses—/ʝ/ “(possession of -----)” and /ʔ/ “(possession of -----)”.

4–8 Hebrew and English verb semantic translation

The English gloss “comes” is present in the top two drawings and the gloss “goes” in the bottom two. However the Hebrew gloss (represented

phonemically) /bá:/ is present in the top left and bottom right drawings with /ya:sá:/ for the top right and bottom left drawings. Examine the two top drawings and figure out why it is appropriate for an English speaker to use the same morpheme for both of these different situations. You should see that in both the figure is moving toward the speaker. Similarly in the bottom two drawings the figure is moving away from the speaker. Now examine the top left and bottom right drawings and figure out what these two situations have in common such that a Hebrew speaker would use /bá:/. In both a figure is entering a house (or enclosure). Similarly in the top right and bottom left drawings the figure is leaving an house (or enclosure) and /ya:sá:/ is used. All four situations involve a person moving, but English and Hebrew classify (and label with morphemes) these actions differently. In Hebrew, regardless of how the person is moving relative to a speaker, the movement is classified by whether the movement is into or out of an enclosure. Therefore, in order to be able to translate the Hebrew morphemes /bá:/ and /ya:sá:/ into English (assuming the translation would not wish to use English morphemes such as “exit” or “enter”) it would be necessary to know the position of the speaker relative to the person moving.

4–9 Semantic treatment of English word in a dictionary

A. I will use my paperback Merriam-Webster Dictionary (2004). Let’s look at how it deals with “star.” “star” has two entries, one for a noun (“¹star n”) and the other for a verb (“²star v”). For this problem I will limit just to the noun listing (but the three verb sub-headings deal only with the noun headings 4 – 6 below). The noun entry has six sub-headings (what the dictionary terms “senses”). Number 1 – 3 refer to “star” as a celestial entity, albeit in different ways—#1 as a luminous body in the sky, #2 and 3 as an astrological guide. To me these seem sufficiently related to be extended meanings for the same morpheme (whichever would be the core meaning). Sub-heading 4 refers to a representational figure such as an asterisk, sub-headings 5 and 6 refer to a person who fills a leading role in a play/movie and is a “brilliant performer.”

B. I believe that the first sub-heading is the denotative meaning in my usage. I am a member of an amateur astronomy group and so this would determine what the core usage would be.

C. I am not sure about the use of the asterisk figure since it is nothing, nor looks nothing, like my core meaning of “star.” However there may be a cultural tradition of associating this mark with the “star” entity. But sub-headings 5 and 6, in my usage, share no semantic relationship with the core and extended meanings discussed above and should have been listed in the dictionary separately.

4–10 Semantic mapping

I will use the domain “club meeting” and let me provide more than the required ten morphemes: agenda, report(s), voting, officer(s), election(s), member(s), advisor(s), gavel, minutes, adjourning, planning—fundraising, events, recruitment (I chose to put these as kinds of planning but they could have been listed separately)

semantic function: nouns—agenda, report(s), officer(s), election(s), member(s), advisor(s), gavel, minutes; verbs—voting, adjourning, planning—event, fund raising, recruitment

(I am unable to think of modifiers distinctive to the domain

relationships—election(s), officer(s), members, and voting occur in combination; election and voting in sequence and officer as the outcome; planning also with voting; agenda and minutes in combination, perhaps also with adjournment

Chapter Five**5–1 Klingon stem and affix identification**

The text suggests five criteria for identifying stems. We will use two of these here (actually only four of the criteria apply since the ability to stand as a free morpheme does not occur in this corpus)—frequency of occurrence and phoneme length. In regard to frequency, two morphemes occur many times—/ʃ/ occurs 8 times and /ɾ/ 5 times; three occur 4 times /wi/, /li/, and /meⁱ/; one occurs three times /da/; and nine occur only once or twice /doʃ/, /yuq^h/, /deʃ/, /duj/, /doʃ/, /nin/, /yaʃ/, /mɛb/, and /puq/. This last group are probably stems. The second criterion is phoneme length. Nine morphemes are three phonemes in length (/doʃ/, /yuq^h/, /deʃ/, /duj/, /doʃ/, /nin/, /yaʃ/, /mɛb/, and /puq/; four are of two phonemes in length /da/, /wi/, /li/, and /meⁱ/, and two are one phoneme /ʃ/ and /ɾ/. Based on this criterion the nine morphemes of three phonemes in length are probably stems, and the two criteria match so these are stems and the rest are affixes, which, based on where they occur in the word relative to the stem makes them suffixes.

5–2 Turkish affix orders

affixes that occur as

order #1

/ ler /

order #2

/ im /

gloss

“-s”

“(first person
possessive pronoun)”

order label

noun number

possessive pronoun
for person

/ in /	“(second person possessive pronoun)”	
<u>order #3</u>		
/ iz /	“(plural for possessive pronoun)”	possessive pronoun number
<u>order #4</u>		
/ den /	“from”	
/ de /	“in”	locative/preposition
/ e /	“to”	
/ li /	“with”	

5-3 English plural affix allomorphs

<u>word</u>	<u>final consonant</u>	<u>suffix</u>	<u>word</u>	<u>final consonant</u>	<u>suffix</u>
“cat”	/ -t /	/ -s /	“lathe”	/ -ð /	/ -z /
dog	/ -g /	/ -z /	bar	/ -j /	/ -z /
lap	/ -p /	/ -s /	catch	/ -č /	/ -əz /
lab	/ -b /	/ -z /	judge	/ -j /	/ -əz /
eye	/ -a ⁱ /	/ -z /	ham	/ -m /	/ -z /
zoo	/ u ^w /	/ -z /	moon	/ -n /	/ -z /
cliff	/ -f /	/ -s /	hall	/ -l /	/ -z /
love	/ -v /	/ -z /	gash	/ -š /	/ -əz /
truck	/ -k /	/ -s /	garage	/ -j /	/ -əz /
myth	/ -θ /	/ -s /	buzz	/ -z /	/ -əz /
feud	/ -d /	/ -z /	song	/ -ŋ /	/ -z /
axe	/ -s /	/ -əz /			

Since there is no difference in meaning among the forms /s/, /z/, and /əz/ (all mean plural) and each occurs in a different phoneme environment, they can be considered to be phonemically conditional allomorphs. Allomorph → /s/ when the stem final is a vl stop or a vl lab-dental or dental fricative; allomorph → /z/ when the stem final is a vd stop, vd lab-dent or dental fric, resonant, or vowel; allomorph → /əz/ when the stem final is an alv fricative or affricate.

5-4 Turkish suffix allomorphs

A) I will list each group of morphemes having the same gloss (and we will assume that there is no issue with this—i.e. the glosses are identical), and thereby consider them as sets of allomorphs

“to”	{ / a / / e / }	“-s”	{ / lar / / ler / }	“from”	{ / ten / / dan / / den / / tan / }
------	--------------------	------	------------------------	--------	--

B) Examining the phonemic environments of these set of morphemes leads to the conclusion that they occur in complementary distribution. The catch here is that the allomorph variation involves two different environments and one of them is not easy to see since it involves the next preceding vowel rather than the contiguous phoneme environment, which is the other environment. The variation is as follows:

suffix vowel → /a/ if next prec. vowel is central or back; suffix vowel → /e/ if next prec. vowel is front; for suffixes with an initial alveolar stop, stop → /t/ if preceding consonant is vl; stop → /d/ if prec. consonant is vd; This would be phonemically conditioned allomorphic variation; “from the birds” would be /kuʃlardan/.

5-5 Comparison of Dinka and English plural morphologically conditioned allomorphs

In the Dinka forms the following phonemic changes occur in stems to change from singular to plural:

- | | |
|------------------------------|--|
| 1. vowel lengthening | 5. change of vowel and vowel lengthening |
| 2. vowel shortening | |
| 3. loss of vowel | 6. change of vowel |
| 4. change of final consonant | |

In English numbers 1-5 do not occur. The one English form that has an allomorphic change similar to Dinka is “man” to “men” a vowel change (from /æ/ to /e/)—similar to the Dinka /gol/ to gal/.

5-6 Word order for English multiple adjectives

		position →						
		6	5	4	3	2	1	house:
In the frame “John bought _____ house:		_____	_____	_____	_____	_____	_____	
position		6	5	4	3	2	1	
a (an)	attractive	blue	brick	large	modern			
the	bad	brown	stone	small	new			
this	good	white	wood(en)	little	old			
his	nice							
my								
your								

5-7 Generative grammar

Let's begin with the assumption that sentence #1 is the base form consisting of an agent (“child”), and action (“love”), and a recipient of the action (“toy”). Thus this base form could be generated by the replacement rules:

- I. $S \rightarrow \text{NounPhrase} + \text{Verb Phrase (usually abbreviated as NP + VP)}$

II. NP $\rightarrow$ Article + Noun

III. VP $\rightarrow$ Verb + NP

(Syntactic rules are “recursive.” This means that your brain does not simply run through the rule sequence only once to generate a sentence. It has the ability to cycle back through earlier replacements and therefore the NP generated in rule III will be expanded into Article + Noun by Rule II.)

Three conditions or markers need to be added in order to complete the rules for sentences #1 and 2 (in other words, to have the rules generate the sentences and ONLY these sentences). First we will need to indicate whether the subject NP’s Noun is singular or plural and second whether the Verb is in the present or future tense. We will rewrite rules II and III as:

$$\text{II. NP} \rightarrow \text{Article} + \text{Noun} + \left\{ \begin{array}{l} \text{Singular} \\ \text{Plural} \end{array} \right\}$$

$$\text{III. VP} \rightarrow \text{Verb} + \left\{ \begin{array}{l} \text{Present} \\ \text{Future} \end{array} \right\} + \text{NP}$$

Our replacement rules are now:

I. S $\rightarrow$ NP + VP

$$\text{II. NP} \rightarrow \text{Article} + \text{Noun} + \left\{ \begin{array}{l} \text{Singular} \\ \text{Plural} \end{array} \right\}$$

$$\text{III. VP} \rightarrow \text{Verb} + \left\{ \begin{array}{l} \text{Present} \\ \text{Future} \end{array} \right\}$$

The curved brackets represent alternatives. The reason for these conditions is that (in English) in the present tense a singular subject will have to have a “-s” on the verb, but not in the future tense. Running through these twice (selecting only the Singular for Noun) will generate:

a. Article + Noun + Singular + Verb + Present + Article + Noun + Singular

b. Article + Noun + Singular + Verb + Future + Article + Noun + Singular

Now we must add the third condition/marker. Some nouns will be “agents,” capable of doing action, such as “child”, and others will be “recipients” that have action done to them, such as “toy”. We must make this condition so that we do not generate a sentence such as: *The toy loves the child. The agent noun will occur in the subject part of the

sentence and the recipient noun in the object part. We make sure of this by employing a rule that is context specific, i.e. it only operate on a particular string of elements. We can now take care of the existence of the correct nouns as subject or object, and the presence of the “-s” suffix on the verb for a singular subject in the present tense. We will use brackets to indicate the specific contextual elements needed for the rule to operate:

IV. [Article + Noun + Singular + Verb + Present] $\rightarrow$ Article +
Noun_{agent} + Singular + Verb + “-s”

V. [Verb + $\left\{ \begin{array}{c} \text{Present} \\ \text{Future} \end{array} \right\}$ + Article + Noun] $\rightarrow$ Verb + $\left\{ \begin{array}{c} \text{Present} \\ \text{Future} \end{array} \right\}$
+ Article + Noun_{recipient}

And for the Future tense:

VI. [Verb + Future] $\rightarrow$ “will” + Verb

Rules #I – VI will be able to generate the appropriate structure (we haven’t supplied a lexicon yet) for sentences #1 and 2.

For the passive sentences #3 and 4 we will need two more contextual rules that are labeled transformations since they operate on a base string of syntactic elements and rearrange or transform them: (transformations are optional rules)

T₁. [Article + Noun_{agent} + Singular + Verb + Present + “-s” + Article +
Noun_{recipient}] $\rightarrow$ Article + Noun_{recipient} + “is” + Verb + “-ed”
+ “by” + Article + Noun_{agent}

T₂. [Article + Noun_{agent} + “will” + Verb + Article + Noun_{recipient}] $\rightarrow$
Article + Noun_{recipient} + “will” + “be” + Verb + “-ed” + “by” +
Article + Noun_{agent}

Finally the Lexicon which supplies the actual words (note that for simplicity’s sake we simply inserted some words into our rules above)

Noun_{agent} $\rightarrow$ child Noun_{recipient} $\rightarrow$ toy Verb $\rightarrow$ love

These six syntactic rules and the two transformations will generate our four sentences and only these sentences.

INDEX

- acoustic phonetics, 15
- adjective, 92
- adverb, 92
- affix, 102 *ff*, 114 *ff*, 129
- affix order, 105
- affixes, obligatory, 105
- affixes, optional, 105
- affrication (cf. consonant co-articulation)
- affricate, 28, 39, 52, 57
- airstream, 15, 19 *ff*, 77
- allomorph, 107 *ff*
- allomorph, language history, 110
- allophone, 59 *ff*, 67
- allophonic consistency, 64
- allophonic conditioning, 64
- allophonic variation, 57, 73
- alphabetic writing system, 76
- alveolar (cf. places of consonant articulation)
- alveolar ridge, 20
- alveo-palatal (cf. places of consonant articulation)
- analogy, 96, 99
- anthropological linguistics, 9
- anthropology and linguistics, 9
- antonymy, 98
- arbitrariness (of speech signal significance), 78
- articulation, 18 *ff*
- articulatory configuration, 24
- articulatory phonetics, 17, 18 *ff*
- aspect (cf. semantic relationships)
- aspiration (cf. consonant co-articulation), 27
- assimilation (allophonic variation), 65
- associative (cf. semantic relationships)
- associative (cf. semantic relationships)
- back (cf. tongue)
- back (cf. vowel tongue positions)
- back glide (cf. glide), 32
- bilabial (cf. places of consonant articulation)
- causative (cf. semantic relationships)
- central glide (cf. glide), 32
- central (df. vowel tongue positions)
- cognitive skills, 2
- comparison (cf. semantic relationship)
- complementary distribution, 60, 118
- complex vowel (cf. glide)
- componential analysis, 86
- compound stems, 117
- concord, 118
- conditional (cf. semantic relationships)
- conjunctive (cf. semantic relationships)
- connotative meaning, 89
- consonant, 19, 23, 29, 71
- consonant chart, 24, 29, 38
- consonant clusters, 72
- consonant co-articulations, 26 *ff*.
- constituent, 119
- construction, 119
- co-occurrence (of phoneme sequence and gloss), 82, 90
- co-occurrence pattern (cf. domain)
- copulative (cf. semantic relationships)
- core meaning, 89
- corpus, 55, 56, 81
- denotative meaning, 89
- dental (cf. places of consonant articulation)
- derivation, 117
- descriptive linguistics, 10

- developmental linguistics, 10, 12
- diacritic marks, 37, 39
- dictionary model, 94
- distinctive features, 51 *ff.*
- distributional analysis, 59 *ff.*
- distributional analysis, strategy, 63
- domain, 94 *ff.*
- egressive (cf. airstream)
- embedded meanings, 86
- English monosyllable formula, 80
- exaggeration (cf. semantic relationships)
- exercises for Chapters Two–Five, 136 *ff.*
- exercises, suggested answers, 156 *ff.*
- extensional meaning, 89
- forward glide (cf. glide)
- frame, 111, 123, 126
- free morphemes, 113
- frequency test (cf. stems and affixes), 103
- fricative (cf. manner of articulation)
- generative grammar, 131
- glide, 31 *ff.*
- gloss (corpus), 56, 81, 86
- glottal (cf. places of consonant articulation)
- glottalization (cf. consonant co-articulations)
- glottis (cf. vocal cords)
- grammar, 5, 6, 100 *ff.*
- graph (notation), 36
- high (cf. vowel tongue positions)
- historical linguistics, 10, 11
- homophones, 90, 116
- human evolution, 7
- infix, 104
- inflection, 117
- ingressive (cf. airstream)
- item, 56
- International Phonetic Alphabet, 18, 131
- intransitivity (cf. semantic relationships)
- knowledge skills (cf. cognitive skills)
- labial (cf. places of consonant articulation)
- labialization (cf. consonant co-articulations)
- labiodental (cf. places of articulation)
- laboratory phonetics, 17
- language, 6
- laryngeal (cf. glottal)
- larynx, 22, 24
- lateral resonant, 28
- length, 33
- lexicon, 93, 132
- linguistic anthropology, 9
- linguistics, 9–12
- locative (cf. semantic relationships)
- low (cf. vowel tongue positions)
- manner of articulation, 18 *ff.*, 24
- meaning (cf. significance, componential analysis), 75 *ff.*
- median resonant, 28
- mental skills (cf. cognitive skills)
- metaphor, 98
- mid (cf. vowel tongue positions)
- minimal pair, 62
- minimal pair analysis, 56
- morpheme, 82 *ff.*
- morpheme (identification) 81 *ff.*
- morphemic analysis, 80 *ff.*
- morpheme significance, 84 *ff.*
- morphologically distributed allomorphs, 109 *ff.*
- morphology, 5, 101 *ff.*
- morphophonemic adjustment, 109
- nasal cavity, 21, 28, 30
- nasal resonant, 28
- nasalization (cf. vowel co-articulations)
- negation (cf. semantic relationships)
- notation, 35 *ff.*
- noun (cf. semantic categories)
- Nuer, 92
- object (cf. semantic relationships)
- openness, 8, 79
- oral cavity, 23, 24, 27, 31, 33, 36
- palatal (cf. places of consonant articulation)
- part of speech (cf. word class)
- pharynx, 21
- phones (speech sounds), 59–65

- phoneme, 46 *ff.*, 51, 67
- phoneme graphs, 67
- phoneme sequences, 69 *ff.*
- phoneme systems, general characteristics, 55
- phonemes of U. S. English, 51 *ff.*
- phonemic analysis (cf. distributional analysis), 55 *ff.*
- phonemic consistency, 66
- phonemic description, 67
- phonemic notation, 69
- phonemics, 47
- phones, 59–65
- phonetic description, 35
- phonetic graph, 41
- phonetic notation, 37 *ff.*
- phonetician, 35
- phonetics, 2–3, 6, 14 *ff.*
- phonologically distributed allomorphs, 108 *ff.*
- phonology, 3, 6, 46 *ff.*
- phonology with four phonemes, 52
- phonology with 150 phonemes, 55
- phrase, 83, 121
- physiological skills, 2
- pitch (cf. tone, prosody)
- places of consonant articulation, 20 *ff.*, 3, 24, 29
- possessive (cf. semantic relationships)
- pragmatics, 4, 6
- prefix, 104
- primate communication, 8
- primate taxonomy, 96
- productivity, 8, 79
- prosodic notation, 43
- prosody, 41, 112
- pulmonic, 19
- reciprocal (cf. semantic relationships)
- recording a corpus, 61
- relational (cf. semantic relationships)
- resonant (cf. manner of articulation)
- root, 117
- root (cf. tongue)
- rounding (cf. vowel co-articulations)
- semantic categories, 92
- semantic relationships, 93
- semantic relationships among affixes (cf. affix)
- semantic “weight” or contribution (cf. stem)
- semantics, 4, 6, 86 *ff.*
- sentence, 83, 111
- sentence, prosodic features, 112
- sentence, semantic structure, 111
- sequential pattern (cf. domain)
- significance, 76 *ff.*
- sociolinguistics, 5, 12
- spectrograph, 15
- spectrogram, 16
- speech, 14
- speech as sound, 15
- speech varieties, 4, 12
- stem, 102 *ff.*
- stop (cf. manner of articulation)
- stress, 41
- subject (semantic relationships), 93
- suffix, 104
- suspicious pairs of sounds, 59, 63
- syllable, 41, 70 *ff.*
- syllable peak (vowel), 71
- syllable structure, 77 *ff.*
- symbols, 79
- synonymy, 98
- syntactic rules, 121 *ff.*
- syntagmatic, 116
- syntax, 5, 111 *ff.*
- taxonomic pattern (cf. domain)
- taxonomy, 96
- tip (cf. tongue)
- tone, 41
- tongue—tip, front, back, root, 21
- tongue height (cf. vowel tongue positions)
- transformation rule, 133
- transitivity (cf. semantic relationships)
- unaspirated stops (cf. consonant co-articulations)
- unrounded vowel (cf. vowel co-articulations)
- uvula, 20, 23

- uvular (cf. places of consonant articulation)
- velar (cf. places of consonant articulation)
- velum, 20
- verb (cf. semantic categories)
- vocal cords, 20, 26
- vocal tract, 3, 15, 20 *ff.*
- vocal tract, cross-section, 21
- voice (cf. semantic relationships)
- voicing (cf. consonant co-articulation)
- voiced (cf. voicing)
- voiceless (cf. voicing)
- vowel, 19, 24, 30, 31, 33, 71
- vowel chart, 25, 31, 33, 39
- vowel co-articulations, 30 *ff.*
- vowel tongue positions—high, mid, low, front, central, back, 24
- word, 83, 101
- word class, 114 *ff.*
- zero morpheme, 106